

Appendix B: Power and sensitivity analysis

Part 1: Power analysis

We conducted an a priori power analysis to determine the required sample size. Since research on the effects of censorship claims and pseudoscientific content on misinformation beliefs is relatively scarce, we grounded our estimates in the broader field of misinformation studies. One meta-analysis (Schmid & Bauer, 2025) showed that exposure to health-related misinformation increased false beliefs ($g = 0.24$). Additionally, O’Brien et al. (2021) demonstrated that misleading claims that embedded scientific content had small-to-medium effects on individuals’ belief perceptions. Based on these findings and large-scale field interventions reported by Roozenbeek et al. (2022), we adopted conservative effect sizes of $f = 0.15$ and $f = 0.20$.

We employed G*Power to conduct an ANOVA (fixed effects, special, main effects, and interactions) to calculate the required sample size given α (0.05), power, and effect size. We used two different power thresholds: 80% (conventional standard) and 95% (conservative standard), with $\alpha = 0.05$. We used the most conservative setting, with a numerator df of 8 (9 conditions - 1).

The power analysis yielded the following required sample sizes:

- Effect size $f = 0.15$, 95% power: $N = 1,019$
- Effect size $f = 0.15$, 80% power: $N = 676$
- Effect size $f = 0.20$, 95% power: $N = 577$
- Effect size $f = 0.20$, 80% power: $N = 384$

Our final analytical sample size is 900, which meets three of the four benchmarks above. We are therefore confident that our sample size provides adequate power to detect the manipulation effects of interest. However, for interaction effects, please see the following sections (Part 2 and Part 3 in Appendix B) below.

Part 2: Sensitivity analysis

We conducted a post hoc sensitivity analysis to determine the minimum three-way interaction effect detectable at several levels of statistical power. Each issue-specific moderation model included 397 participants, seven predictors including the one three-way interaction term, and 389 residual degrees of freedom. With a nominal two-sided alpha of .05, the models had 80%, 75%, and 70% power to detect incremental effects of Cohen’s $f^2 = .0199$, .0176, and .0156, respectively. The statistically significant three-way interactions had Cohen’s f^2 values ranging from .0161 to .0180 (see Table B1). Thus, the observed effects were above the 70% sensitivity threshold; the two scientific-populism interactions were also slightly above the 75% threshold, but none reached the 80% power threshold. Although these interactions were statistically significant in the present sample, effects of this magnitude may not be detected consistently across replications. The findings should therefore be interpreted cautiously and replicated in studies specifically powered to detect higher-order interactions for more conclusive findings on these individual difference moderators.

Table B1. Sensitivity results for three-way interaction models.

Outcome	Moderator	<i>N</i>	<i>b</i>	<i>SE</i>	<i>t</i>	<i>p</i>	Adjusted <i>R</i> ²	Cohen's <i>f</i> ²
Vaccine risk perception	Subjective numeracy	397	−0.661	0.264	−2.501	0.013	0.018	0.0161
Vaccination intention	Subjective numeracy	397	0.088	0.292	0.302	0.763	0.028	0.0002
GMO risk perception	Subjective numeracy	397	−0.225	0.264	−0.851	0.395	−.002	0.0019
GMO consumption intention	Subjective numeracy	397	0.068	0.294	0.23	0.818	0.01	0.0001
Vaccine risk perception	Scientific populism	397	0.485	0.315	1.539	0.125	0.119	0.0061
Vaccination intention	Scientific populism	397	−0.979	0.37	−2.645	0.008	0.014	0.018
GMO risk perception	Scientific populism	397	−0.477	0.304	−1.572	0.117	0.041	0.0063
GMO consumption intention	Scientific populism	397	0.894	0.338	2.642	0.009	0.053	0.0179

*Note: Each row reports the three-way interaction among pseudoscientific content, censorship accusations, and the specified moderator. The *b*-values are unstandardized regression coefficients, and the *p*-values are based on two-sided tests. Cohen's *f*² represents the incremental contribution of the focal three-way interaction and was calculated as t^2 divided by the residual degrees of freedom.*

Part 3: False discovery rate adjustment

To address multiple testing across the higher-order interaction models, we applied the Benjamini–Hochberg (BH) false discovery rate procedure within three conceptually distinct (as pre-registered) outcome families: credibility, vaccine-related outcomes, and GMO-related outcomes. The credibility family included four three-way interaction tests, one for each moderator. The vaccine family included eight tests across vaccine risk and vaccination intention, and the GMO family included eight tests across GMO risk and GMO consumption intention. The moderators were subjective numeracy (main paper), scientific populism (main paper), trust in science (appendix result), and need for cognition (appendix result). This family structure reflects differences in the outcome constructs and analytical samples: credibility was pooled across the vaccine and GMO topics, whereas the vaccine and GMO outcomes were measured in separate topic-specific subsamples.

The BH procedure ranks the unadjusted *p*-values within each family and adjusts them according to their rank and the number of tests in that family (Benjamini & Hochberg, 1995). The resulting BH-adjusted *p*-value, referred to here as the *q* value, indicates whether a result remains significant while controlling the expected proportion of false discoveries among the results declared significant within that family. Thus, $q < .05$ controls the false discovery rate at 5%; it does not mean that each significant result has a 5% probability of being false.

In addition to the BH adjustment, we also tested for a more conservative Bonferroni adjustment. Bonferroni controls the family-wise probability of making one or more Type I errors and hence is possible

to lead to false negatives as well. Thus, we present results under both adjustment procedures to assess the robustness of the findings.

Results for both adjustment robustness checks are in Table B2: Table B2 reports the unadjusted p -values, within-family BH-adjusted p -values (q -values), and Bonferroni-adjusted p values.

BH Results: Of the 20 exploratory three-way interactions, two of the interactions reported in the main text remained significant after the within-family BH-FDR adjustment. These were the subjective numeracy interaction for vaccine risk and the scientific populism interaction for vaccination intention. The scientific populism interaction for GMO consumption intention was significant before adjustment but did not remain significant after BH-FDR adjustment ($p = .009$, $q = .069$).

Bonferroni Results: None of the interactions remained significant under the more conservative within-family Bonferroni adjustment. We treat the BH-FDR results as the primary multiplicity analysis and present the Bonferroni results as a stricter sensitivity analysis. Collectively, these results support cautious interpretation of the higher-order interactions.

Table B2. Three-way interaction models with within-family Benjamini-Hochberg FDR and Bonferroni adjustments.

Family	Moderator	Outcome	b	SE	t	p	BH-FDR q	Bonferroni p
Credibility	Subjective numeracy	Credibility	-0.247	0.157	-1.573	0.116	0.411	0.464
Vaccine outcomes	Subjective numeracy	Vaccine risk	-0.661	0.264	-2.501	0.013	0.034	0.102
Vaccine outcomes	Subjective numeracy	Vaccination intention	0.088	0.292	0.302	0.763	0.763	1
GMO outcomes	Subjective numeracy	GMO risk	-0.225	0.264	-0.851	0.395	0.632	1
GMO outcomes	Subjective numeracy	GMO consumption intention	0.068	0.294	0.23	0.818	0.858	1
Credibility	Scientific populism	Credibility	-0.018	0.181	-0.1	0.92	0.949	1
Vaccine outcomes	Scientific populism	Vaccine risk	0.485	0.315	1.539	0.125	0.216	0.997
Vaccine outcomes	Scientific populism	Vaccination intention	-0.979	0.37	-2.645	0.008	0.034	0.068
GMO outcomes	Scientific populism	GMO risk	-0.477	0.304	-1.572	0.117	0.467	0.935
GMO outcomes	Scientific populism	GMO consumption intention	0.894	0.338	2.642	0.009	0.069	0.069
Credibility	Trust in science	Credibility	-0.188	0.149	-1.267	0.206	0.411	0.823
Vaccine outcomes	Trust in science	Vaccine risk	-0.397	0.265	-1.498	0.135	0.216	1
Vaccine outcomes	Trust in science	Vaccination intention	0.214	0.289	0.742	0.458	0.611	1
GMO outcomes	Trust in science	GMO risk	-0.284	0.246	-1.156	0.248	0.497	1

Family	Moderator	Outcome	<i>b</i>	<i>SE</i>	<i>t</i>	<i>p</i>	BH-FDR <i>q</i>	Bonferroni <i>p</i>
GMO outcomes	Trust in science	GMO consumption intention	0.358	0.275	1.303	0.193	0.497	1
Credibility	Need for cognition	Credibility	-0.011	0.172	-0.064	0.949	0.949	1
Vaccine outcomes	Need for cognition	Vaccine risk	-0.81	0.299	-2.708	0.007	0.034	0.056
Vaccine outcomes	Need for cognition	Vaccination intention	0.156	0.336	0.463	0.643	0.735	1
GMO outcomes	Need for cognition	GMO risk	0.093	0.284	0.328	0.743	0.858	1
GMO outcomes	Need for cognition	GMO consumption intention	-0.057	0.32	-0.18	0.858	0.858	1

*Note: Each row reports the three-way interaction among pseudoscientific content, censorship accusations, and the specified moderator. The credibility models pooled the vaccine and GMO topics and controlled for topic; the vaccine and GMO models were estimated in their respective topic-specific experimental subsamples. The *b*-values are unstandardized regression coefficients, and the *p*-values are based on two-sided tests. BH-adjusted *p*-values, referred to as *q*-values, were calculated separately within the credibility family (four tests), vaccine-outcome family (eight tests), and GMO-outcome family (eight tests). A value of $q < .05$ indicates significance under a 5% false discovery rate within the corresponding family. Bonferroni-adjusted *p*-values were calculated within the same families and provide a more conservative test controlling the familywise error rate.*