

Supplemental Material for

Targeting Audiences’ Moral Values Shapes Misinformation Sharing

1 Stimuli

Study 1 & Study 2

Table S1

List of headlines selected for Study 1 and Study 2

#	Headline	True/False	Sentiment
1	Man Shoots Off His Left Ear Taking Selfies With Gun	False	Negative
2	Refugees Have 100 Times Greater Rate Of Tuberculosis Than National Average	False	Negative
3	Starbucks Is Giving Out Free Lifetime Passes On Its 44th Anniversary	False	Negative
4	Man Infects 586 People With HIV On Purpose, Plans On Infecting 2,000 More Before 2024	False	Negative
5	Trillionaires Now Exist!	False	Negative
6	America Is Now Reducing CO2 Emissions Much Faster Than Other Developed Countries	False	Positive
7	Black And White Wealth Gap Is Closing Fast	False	Positive
8	John Travolta Takes A New Wife After The Death Of Kelly Preston	False	Positive
9	Polar Bears Are Strong and Healthy Across Alaska	False	Positive
10	Taco Bell Reportedly Going Out of Business	False	Positive
11	Twitter Is Making A Dislike Button	True	Negative
12	Crocs is Giving Away Free Footwear to Healthcare Workers	True	Negative
13	Portland Named a New Bridge After The Simpsons Ned Flanders	True	Negative
14	State Department Charging Americans \$2k For Flights Out Of Afghanistan	True	Negative
15	Whitest-Ever Paint Could Help Cool Heating Earth	True	Negative
16	Hole In The Ozone Layer Will Totally Heal Within 50 Years	True	Positive
17	An Invisible Sculpture Sold for \$18K	True	Positive
18	A California Man Sued A Psychic For Allegedly Failing To Remove A Curse	True	Positive
19	Bluetooth Technology Was Named After A Viking King	True	Positive
20	Scientists Detect Cocaine In Freshwater Shrimp	True	Positive

Note. Headlines within each combination of veracity and sentiment are ordered by the criterion described in the Results section of Study 1a.

Table S2*Classification of headline moral framing using a fine-tuned classifier*

#	Headline	Classifier
1	Man Shoots Off His Left Ear Taking Selfies With Gun	Neutral
2	Refugees Have 100 Times Greater Rate Of Tuberculosis Than National Average	Neutral
3	Starbucks Is Giving Out Free Lifetime Passes On Its 44th Anniversary	Neutral
4	Man Infects 586 People With HIV On Purpose, Plans On Infecting 2,000 More Before 2024	Neutral
5	Trillionaires Now Exist!	Binding
6	America Is Now Reducing CO2 Emissions Much Faster Than Other Developed Countries	Neutral
7	Black And White Wealth Gap Is Closing Fast	Neutral
8	John Travolta Takes A New Wife After The Death Of Kelly Preston	Neutral
9	Polar Bears Are Strong and Healthy Across Alaska	Neutral
10	Taco Bell Reportedly Going Out of Business	Neutral
11	Twitter Is Making A Dislike Button	Neutral
12	Crocs is Giving Away Free Footwear to Healthcare Workers	Neutral
13	Portland Named a New Bridge After The Simpsons Ned Flanders	Neutral
14	State Department Charging Americans \$2k For Flights Out Of Afghanistan	Neutral
15	Whitest-Ever Paint Could Help Cool Heating Earth	Neutral
16	Hole In The Ozone Layer Will Totally Heal Within 50 Years	Neutral
17	An Invisible Sculpture Sold for \$18K	Neutral
18	A California Man Sued A Psychic For Allegedly Failing To Remove A Curse	Neutral
19	Bluetooth Technology Was Named After A Viking King	Neutral
20	Scientists Detect Cocaine In Freshwater Shrimp	Neutral

Note. Headlines were classified using a fine-tuned BERT classifier as described in the methods of Study 1a and Study 3. The classifier was not specifically trained for this task and did not see any of the headlines and comments at any point. Classification shows that the headlines have indeed, with only a single exception, a morally neutral framing (neither Binding nor Individualizing) and the moral sentiment is thus solely expressed in the accompanying comments.

Table S3*Classification of comments' moral framing using a fine-tuned classifier - Part I*

#	Comment	Framing	Classifier
1	Wow, I love this! Artists sticking it to the gullible rich! Take their money and use it for something more worthwhile, like healthcare!	Ind.	Neutral
2	Wow! I love a good rags to riches entrepreneurial story! In America any dream, no matter how wild, can come true.	Bind.	Neutral
3	Wow! This is such a creative prank, and very fun to see. It really made me laugh. I wish I were that creative myself.	Neutral	Neutral
4	I read this article and can't believe it. We have so many civil rights leaders who go nameless and we give it to another white man!	Ind.	Ind.
5	I read this article and I can't believe it! This bridge should be named after a great American patriot, not a cartoon character!	Bind.	Bind.
6	I read this article and am surprised a bridge named after a Simpsons character?! Ridiculous! People have too much time on their hands!	Neutral	Neutral
7	Horrible that our healthcare workers are not taken care of by our government and have to rely on handouts from a private company.	Ind.	Ind.
8	What a waste of money. Companies like this should consider donating their products to the men and women of our armed forces who are protecting this country right now.	Bind.	Neutral
9	Man, Crocs are so ugly. You couldnt pay me to wear them. No idea why anyone would wear them, free or not these things are more pool floaty than shoes!	Neutral	Neutral
10	What a waste of money. They should have used that money to help people who actually need it, like their underpaid employees!	Ind.	Ind.
11	What a waste! They should have given these to our military personnel who sacrifice their lives for our country.	Bind.	Bind.
12	What a stupid PR stunt! They could have done much more useful things instead! Give me a break Starbucks.	Neutral	Neutral
13	Im glad this was caught by these scientists! We must care for our nature & wildlife - science for the win! Its only fair for humans to correct the harm we caused.	Ind.	Ind.
14	Im glad this was caught! We have to make sure our countrys nature is free of these impurities! The less of of these disgusting drugs we have in our environment the better!	Bind.	Bind.
15	Wow! I wonder how the cocaine affected these shrimp! Just another example of the butterfly effect I guess. The interdependence of the ecosystem is fascinating.	Neutral	Neutral

Note. Comments were classified according to the description in Study 1a Classification shows that the detected comments' moral sentiment aligns in all but few cases where the classifier detects no moral sentiment.

Table S4*Classification of comment moral framing using a fine-tuned classifier - Part II*

#	Comment	Framing	Classifier
16	Good to see that he found happiness again. I think its important to find someone you truly care for! No one should suffer forever.	Ind.	Ind.
17	Good to see that he found happiness again. Clearly, he values the sanctity of marriage over a frivolous single life-style.	Bind.	Bind.
18	Nice to see that he found happiness again! Always liked this guys movies. Saturday Night Fever is a classic!	Neutral	Neutral
19	Unbelievable! We have people suffering on the streets and growing child poverty yet those at the top are stealing more and more money away from society. Boo!	Ind.	Ind.
20	Just another hit piece against those geniuses in our society that work tirelessly to build wealth and products that help all of our lives! You go Jeff Bezos, screw you to the haters!	Bind.	Bind.
21	This is so stupid. Cant believe this is the type of stuff we are spending our time writing and reading articles about. Who cares if trillionaires exist?	Neutral	Neutral
22	Another company exploiting workers with unlivable wages is going bustthats a good thing! One more step toward achieving equity! Down with corporations that mistreat workers.	Ind.	Ind.
23	Our bodies are temples! Im glad that we no longer have to wonder what type of impure chemicals were consuming with Taco Bell. Good riddance to this toxic, disgusting food!	Bind.	Bind.
24	Wow, FINALLY Taco Bell is going out of business! I never liked their food anyway, and there are so many other good alternatives out there for fast food. Good riddance!	Neutral	Neutral
25	Just another fad for the rich. Yes, greenhouse gas emissions are important but this only helps the image of the wealthy when there are hundreds of people dying of hunger everyday.	Ind.	Ind.
26	This paint is not going to make our country better in any way. The paint my parents used to paint their house was fine, they didnt need this new techy paint which is probably toxic anyways.	Bind.	Neutral
27	Literally what is the point? Probably completely impossible to manufacture this stuff at scale anyways. Also, would be dirty in 5 seconds. and no longer white.	Neutral	Neutral
28	I read this article and I just think that this will be another tool to hound and harrass women and minorities on Twitter!	Ind.	Ind.
29	I read this article and I just think that this will be another tool to cancel patriots and smear proud Americans on Twitter!	Bind.	Bind.
30	I read this article and I'm not sure I like this. Pun intended. Do we really need another button? This is not reddit!	Neutral	Neutral

Note. Comments were classified according to the description in Study 1a Classification shows that the detected comments' moral sentiment aligns in all but few cases were the classifier detects no moral sentiment.

Table S5*Classification of comment moral framing using a fine-tuned classifier - Part III*

#	Comment	Framing	Classifier
31	It is unacceptable that people are not helped in a situation like this. It is not fair that only the rich are getting rescued!	Ind.	Ind.
32	Disgusting! I cannot believe that our country would treat patriots risking their lives for our country like this! This is treason!	Bind.	Bind.
33	I didnt know thats how evacuation flights workand I dont think its a good thing.	Neutral	Neutral
34	This is wild! Its good someone is suing (even for this reason) to stop psychics. They prey on vulnerable people and cause so much harm!	Ind.	Ind.
35	This is wild! Its good someone is suing (even for this reason) to stop psychics. Psychic practice is not Christian, its not American, and its not right!	Bind.	Neutral
36	This is wild! Good for him to try to get his money back. The lengths people need to go to is unbelievable.	Neutral	Neutral
37	Bluetooth is named after a Viking king who unified his people, bringing peace and equality to a normally violent group. I always loved how viking women had power alongside the men. Thinking of that amazing legacy will make me smile every time I pair my earbuds to my phone	Ind.	Ind.
38	Bluetooth is named after a brave Viking king who unified his countrymen and converted them to Christianity. Thinking of that amazing legacy will make me smile every time I pair my earbuds to my phone	Bind.	Bind.
39	Bluetooth is named after a brave Viking king. The naming of companies is usually so boring but this actually makes sense and is cool!	Neutral	Neutral
40	This goes to show that we can make the world a better place with health and justice for allif we only try together as one world!	Ind.	Ind.
41	This goes to show that our natural world always finds a way to heal- especially when us Americans lead the way to keep our air clean and pure!	Bind.	Bind.
42	This is good newsone less thing to worry about! Maybe Ill need less sunscreen next time I go to Australia!	Neutral	Neutral
43	This type of tragedy is preventable! When we flood our communities with guns, harm like this is inevitable. GUN-CONTROL NOW!	Ind.	Ind.
44	Wow, this type of show-boating mistake is shameful. The reason our great country allows guns is to preserve our freedom. This kind of idiocy defiles our countrys values. Follow firearm protocols!	Bind.	Neutral
45	Wow! Incompetence like this is dangerous! It is hard to believe people can be so negligent. I am sure he will regret this mistake! YIKES!	Neutral	Neutral

Note. Comments were classified according to the description in Study 1a Classification shows that the detected comments' moral sentiment aligns in all but few cases where the classifier detects no moral sentiment.

Table S6*Classification of comment moral framing using a fine-tuned classifier - Part IV*

#	Comment	Framing	Classifier
46	This is insane to see the amount of harm one person can do. And of course, once again, its not the rich who can afford HIV drugs that are targeted. So unfair to these victims.	Ind.	Ind.
47	So disgusting to see people infected with this nasty disease on purpose. This is absolutely criminal. There has to be a way we can protect our community against this kind of attack.	Bind.	Bind.
48	Wow this is so horrible! What kind of world do we live in where someone walks around doing something so mean and terrible. Ugh, theres so much bad news lately!	Neutral	Neutral
49	This is so unjust! Its outrageous that so many people across the world do not have access to healthcare that could prevent treatable conditions like this. We need to take care of these refugees!	Ind.	Ind.
50	This is so disgusting! Tuberculosis was almost gone in America, a testament to our cleanliness and hygiene. We need to protect the American people from foreign diseases!	Bind.	Bind.
51	Oh no, I sure hope I wont catch Tuberculosis now. I already have plenty of things to worry about, sure dont want to have to worry about this as well!	Neutral	Neutral
52	This is amazing news great to see that we can care for the wildlife. Its only fair that polar bears can thrive in their natural habitat!	Ind.	Ind.
53	Alaska is such a great representation of American values rugged terrain and hard work, with strong and thriving polar bears	Bind.	Neutral
54	This is awesome news! I love polar bears. And they are important to the ecosystem of Alaska!	Neutral	Neutral
55	This is great news! Black Americans have been suffering so long as a result of racism and discrimination its great that they are finally getting their fair share!	Ind.	Ind.
56	This is great news! This beautiful country was built on the notion that everyone has a shot to live the American dream! This is good for our country!	Bind.	Bind.
57	Good for them! Its always nice to hear good news when everything can seem bleak and hopeless.	Neutral	Neutral
58	Yeah! For too long have we emitted much more than our fair share of emissions, harming the most vulnerable communities across the world. Its good that we are making progress!	Ind.	Ind.
59	Hoorah! This country is leading the way as it always has! American innovation and determination has accomplished monumental tasks in the past and we are doing the same once again!	Bind.	Bind.
60	I had no ideabut this is definitely a good thing! There is still so much to learn. So many new technologies are needed to get us where we need to be. What a time to be alive!	Neutral	Neutral

Note. Comments were classified according to the description in Study 1a Classification shows that the detected comments' moral sentiment aligns in all but few cases were the classifier detects no moral sentiment.

Study 3

Table S7

Exemplary tweets showcasing moral framing and/or political ideology on COVID vaccinations and mandates

Topic	Description	Example
Nonmoral	Does not contain moral language	They really think mandating vaccines on airplanes is gonna sway the unvaccinated, lol. I guess Im gonna just drive...
Individualizing	Focused on individual rights and well-being	Under 12s are unvaccinated! We need to ensure all primary schools have safe air to prevent mass infections.
Binding	Focused on group preservation	Common law, natural law, God's law. I will never consent and am both disgusted and horrified by people's acquiescence... My body belongs to GOD!
Liberal user (pro-vax)	Endorsing COVID vaccines and mandates	Lets get vaxxed!
Conservative user (anti-vax)	Opposing COVID vaccines and mandates	No. Do Not get Vaccinated!

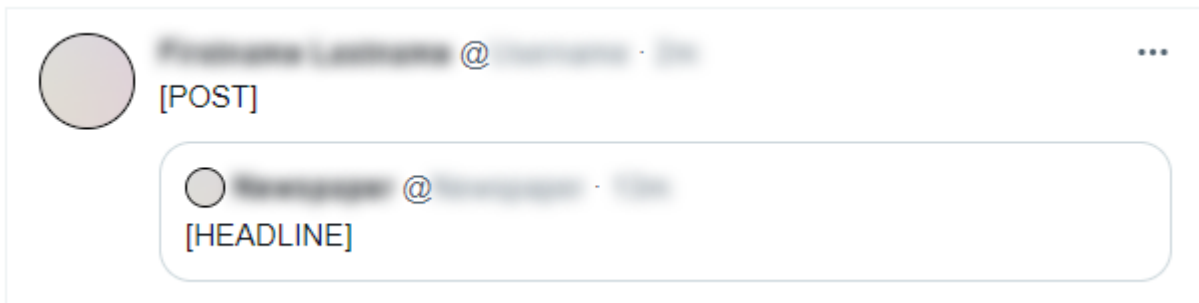
2 Questionnaire Items (Study 1a, 1b, 2)

Introduction

On the next few pages, you will see examples of social media posts that look somewhat like this:

Figure S1

Example stimulus presentation for the introduction



Each example contains a news headline a user has shared (here: [HEADLINE]) and a post the user has written about the news headline (here: [POST]). For this study, we are leaving out some information about the post (for example, who posted it and when). Please answer each of the questions as if you had come across the post while using social media (e.g., Twitter or Facebook). In total, you will answer questions about 5 posts.

Headline-Level

On this page, please focus on the headline the user has shared:

Figure S2
Example stimulus presentation for headline-level items

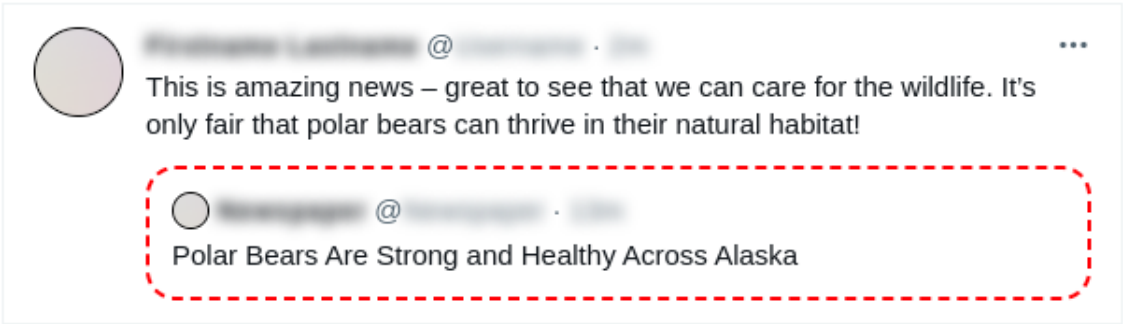


Table S8
In your opinion, the post the user has written about the headline is

Unbelievable	1	2	3	4	5	6	7	Believable
Uncontroversial	1	2	3	4	5	6	7	Controversial
Unsurprising	1	2	3	4	5	6	7	Surprising
Uninteresting	1	2	3	4	5	6	7	Interesting
Negative	1	2	3	4	5	6	7	Positive

Post-Level

On this page, please focus on the post the user has written about the headline:

Figure S3

Example stimulus presentation for post-level items

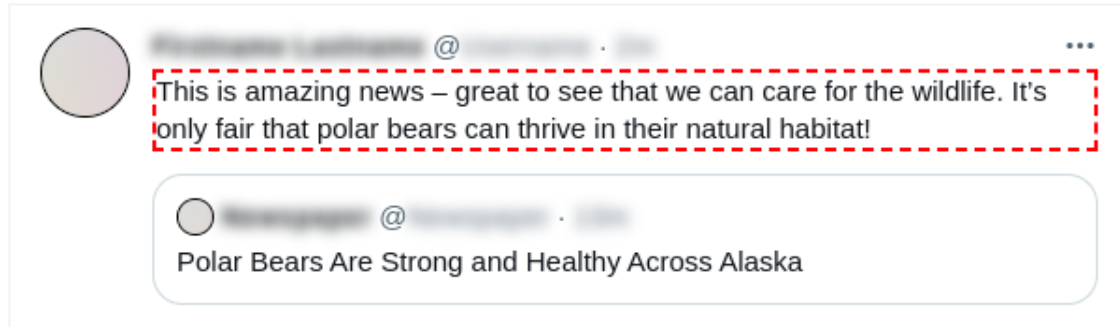


Table S9

In your opinion, the post the user has written about the headline is

Uncontroversial	1	2	3	4	5	6	7	Controversial
Unsurprising	1	2	3	4	5	6	7	Surprising
Uninteresting	1	2	3	4	5	6	7	Interesting
Negative	1	2	3	4	5	6	7	Positive

Table S10

How much do you agree or disagree with the post the user has written about the headline?

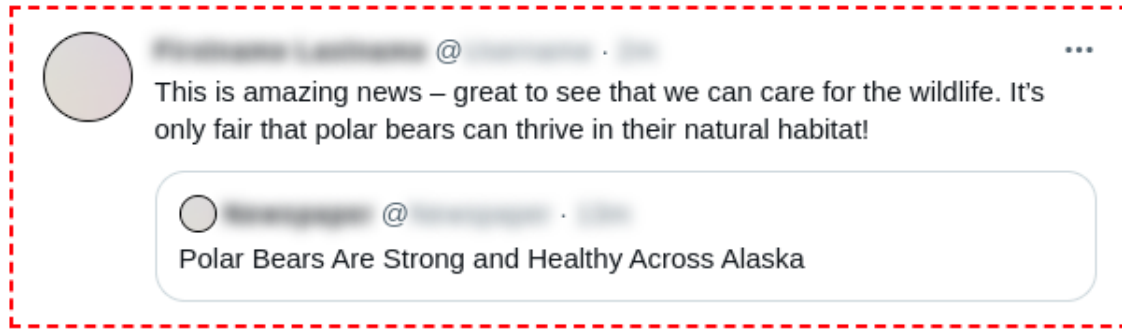
Strongly disagree	1	2	3	4	5	Strongly agree
-------------------	---	---	---	---	---	----------------

Table S11

How much does the post the user has written about the headline align with your values?

Strongly opposed to my values	1	2	3	4	5	6	7	Strongly aligned with my values
-------------------------------	---	---	---	---	---	---	---	---------------------------------

On this page, please focus on the entire social media post:

Figure S4*Example stimulus presentation for whole-stimulus items***Table S12***If you came across this post, how likely would you be to ...*

... share the post publicly on your social media feed (e.g., retweet or share on your Facebook)	1	2	3	4	5
... 'like' the post (e.g., Twitter or Facebook)?	1	2	3	4	5
... share the post privately in a private message, text message, or email?	1	2	3	4	5
... talk about the post or headline in an offline conversation?	1	2	3	4	5

Table S13*Why did you decide to share (or not to share) the previous post? (Only Study 2)*

It just felt right/wrong without much thought	1	2	3	4	5	I carefully considered the information and implications
It matched my gut-feeling and intuitions	1	2	3	4	5	It matched my knowledge and experiences
It made me feel strongly	1	2	3	4	5	It made me very thoughtful
My social circle posts similar/different posts	1	2	3	4	5	I reflected on my own thoughts and opinions

3 Study 3 Robustness Checks

The data in Study 3 was sampled from was a corpus originally collected in an effort to study misinformation, rumors, and conspiracy theories regarding COVID vaccinations and mandates. However, this data was collected using keywords, which might have led to a significant number of truthful content being included in this data. As a robustness check, we determined the average level of exposure to misinformation for users in our dataset as a proxy for the prevalence of misinformation in our data. We used the API previously used to retrieve users' political ideology (Mosleh & Rand, 2022) to retrieve users' exposure score from 0 (no exposure) to 1 (complete exposure), based on the accounts the user follows, the credibility of sources that these accounts share, as well as the amount of previously identified misinformation shared by these accounts. Note that Mosleh and Rand (2022) found a strong correlation between the misinformation-exposure score and sharing of low-credibility news sources ($b = -0.712$, $p < .001$). Thus, a high exposure score in our sample would indicate a high prevalence of shared misinformation in our dataset. We find an average misinformation exposure score of $M = 0.59$. As a comparison, Mosleh and Rand (2022) identified clusters of low-exposed users with an average exposure score of $M = 0.389$ and a highly-exposed cluster with an average exposure score of $M = 0.501$, indicating that our data likely has a high prevalence of misinformation. Additionally, independent of the likely high prevalence of misinformation in our dataset, it could be that the observed effects are driven mainly or even exclusively by the (few) true information that could be in the dataset. As a robustness check, we created another dataset from the original corpus (Muric et al., 2021), this time only collecting tweets that contain URLs that can be verified as misinformation (e.g., from misleading, low-credibility sources)⁸. Specifically, we filtered for messages with URLs that appear in the most current version of the iffy index from <https://iffy.news/index/>. The iffy index is a collection of nearly 2000

⁸ Note that access to the Twitter API has been highly limited prohibiting data collection from scratch. However, access to larger amounts of tweets might be possible through the original authors of the avax twitter corpus.

online Domains that have been verified to spread incorrect, misleading, and biased information. A previous version of this index was used by Muric et al. (2021) in the original publication of their corpus. Finally, we applied a classifier to detect the moral framing in the tweets analogous to the approach used in Study 3. The final dataset contained 54,885 tweets from 23,680 users. Note that for this dataset, we could not determine users’ political ideology because the API for determining Twitter users’ ideology had, at the time of conducting this robustness check, significant rate limits that prevented us from making the necessary API calls in an acceptable time frame. Therefore, we used stance on vaccination as a proxy for ideology, which was provided by the original authors. Our analysis in Section 5 shows that stance is a reasonable proxy for ideology (anti-vax correlating with conservatism and pro-vax correlating with liberalism).

Table S14 compares each model’s out-of-sample prediction accuracy of engagement, captured by retweet count to that of the null model without predictors (M0) and the other models with predictors (M1–M2). We found that Model 2—which included tweet’s moral framing (Binding and Individualizing) and their interactions with the user’s ideology (liberal and conservative)—predicted engagement more accurately than Model 0 ($\Delta_{ELPD} = 368, SE = 26.5, z = 13.9$) and Model 1 ($\Delta_{ELPD} = 25.6, SE = 19.6, z = 1.31$), indicating the relevance of matching moral framing and individuals’ values for the spread of social media messages. The between-group analyses demonstrated, as hypothesized, that tweets’ Individualizing framing predicted more (1.3 times) engagement when posted by liberal users compared to conservative users ($\beta = 0.26, [0.19, 0.32]$). Analogous, posts with Binding framing received significantly more engagement (1.4 times) when posted by conservative versus liberal users ($\beta = 0.31, [0.23, 0.39]$). However, the within-group analyses, did not show that Individualizing framing predicts significantly more engagement than Binding framing for liberal users ($\beta = -0.07, [-0.12, -0.02]$). Additionally, Binding values (vs Individualizing) led to significantly more engagement (1.9 times) for conservatives ($\beta = 0.64, [0.56, 0.72]$). See an overview of effect sizes and confidence intervals

for model M2 in Table S15.

Table S14

Comparison of models estimating engagement (retweet count) as a function of various predictor variables

Model	Predictors	R^2	z		
			M0	M1	M2
M0	Political ideology	0.40	-	-16.1	-13.9
M1	Moral framing, Ideology	0.41	16.1	-	-1.31
M2	Moral framing, Ideology, Interaction	0.41	13.9	1.31	-

Note. R^2 is a Bayesian analogue to the proportion of within-sample variance explained by a model (not considering varying effects). z is the difference in out-of-sample prediction accuracy between two models divided by its standard error ($z = \Delta_{ELPD}/SE$).

Table S15

Effect of post framing and political ideology on engagement (retweet count)

Between-Group Analysis		Within-Group Analysis	
Framing	Aligned vs Misaligned	Ideology	Aligned vs Misaligned
Binding	0.31 [0.23, 0.39]	Liberal	-0.07 [-0.12, -0.02]
Individualizing	0.26 [0.19, 0.32]	Conservative	0.64 [0.56, 0.72]

Note. The left side of the table shows the difference in the effect of moral alignment vs misalignment on sharing intentions for posts with Binding and Individualizing framing (Between-Group Analysis). Posts with Binding framing are aligned with conservative (rather than liberal) users, while posts with Individualizing framing are aligned with liberal (rather than conservative) users. The right side of the table shows the difference in the effect of moral alignment vs misalignment on sharing intentions for conservative and liberal users message framing (Within-Group Analysis). For liberal users, Individualizing framing is aligned, and for conservative users, there should be no prioritization of one value/framing over the other.

Analogous to Table S14, Table S16 compares each model’s out-of-sample prediction accuracy of engagement, captured by favorite count, to that of the null model without predictors (M0) and the other models with predictors (M1–M2). Supporting Hypothesis 1, Model 2—that included tweets’ moral framing (Binding and Individualizing) and their interactions with the user’s ideology (liberal and conservative)—predicted engagement more accurately than Model 0 ($\Delta_{ELPD} = 389$, $SE = 24.8$, $z = 15.7$) and Model 1

($\Delta_{ELPD} = 36.9$, $SE = 16.7$, $z = 2.21$). The between-group analyses demonstrated, as hypothesized, that the tweets' Individualizing framing predicted more (2.5 times) engagement when posted by liberal users compared to conservative users ($\beta = 0.89$, $[0.42, 1.37]$). Similarly, posts with Binding framing elicited significantly more engagement (1.4 times) for conservatives compared to liberal users ($\beta = 0.32$, $[0.22, 0.41]$). However, the within-group analyses showed that Individualizing framing did not predict significantly more engagement compared to Binding framing for liberal users ($\beta = -0.30$, $[-0.36, -0.24]$). We also found that conservative users received more engagement (2.2 times) for posts with Binding compared to Individualizing framing ($\beta = 0.81$, $[0.72, 0.89]$). See an overview of the effect sizes and confidence intervals for model M2 in Table S17.

Table S16

Comparison of models estimating engagement (favourites count) as a function of various predictor variables

Model	Predictors	R^2	z		
			M0	M1	M2
M0	Political ideology	0.47	-	-15.3	-15.7
M1	Moral framing, Ideology	0.47	15.3	-	-2.21
M2	Moral framing, Ideology, Interaction	0.48	15.7	2.21	-

Note. R^2 , here, is its Bayesian counterpart. z is the standardized difference in prediction accuracy between each model ($z = \Delta_{ELPD}/SE$).

Table S17*Effect of post framing and political ideology on engagement (favorite count)*

Between-Group Analysis		Within-Group Analysis	
Framing	Aligned vs Misaligned	Ideology	Aligned vs Misaligned
Binding	0.32 [0.22, 0.41]	Liberal	-0.30 [-0.36, -0.24]
Individualizing	0.19 [0.12, 0.26]	Conservative	0.81 [0.72, 0.89]

Note. The left side of the table shows the difference in the effect of moral alignment vs misalignment on sharing intentions for posts with Binding and Individualizing framing (Between-Group Analysis). Posts with Binding framing are aligned with conservative (rather than liberal) users, while posts with Individualizing framing are aligned with liberal (rather than conservative) users. The right side of the table shows the difference in the effect of moral alignment vs misalignment on sharing intentions for conservative and liberal users message framing (Within-Group Analysis). For liberal users, Individualizing framing is aligned, and for conservative users, there should be no prioritization of one value/framing over the other.

Overall, these findings show that an alignment of moral framing and ideology increases engagement with social media messages and supports that the results of our main analysis in Study 3 are in fact driven by misinformation. Note that we found disparities in our within-group analyses. Mainly, liberals (pro-vax) did not receive more engagement for using Individualizing compared to Binding framing. And conservatives (anti-vax) received more engagement for using Binding compared to Individualizing framing. This could be caused by Binding values being generally more relevant in the vaccination debate. For example, past work has shown that Binding values relate to both pro- and anti-vaccination behavior (Reimer et al., 2022)—e.g., loyalty to the group increases vaccination behavior and spiritual and physical purity reduces vaccination behavior. Similarly, posts with these framings could also be written more convincingly, independent of their framing. In the Twitter data, we cannot control for the specific arguments made by each group (e.g., Binding posts being written more convincingly or interestingly, independent of the poster’s ideology). However, in the experiments of Study 1 and Study 2 we explicitly control for these factors by keeping the underlying arguments across framing conditions constant. In these cases we do observe significant within-group alignment of message-framing and

participant values on sharing.

4 Detailed Model Formulas

Study 1b

Table S18

R Formulas for Bayesian Multilevel Linear Regression Models

Model	R Formula
M0	<code>brm(z_post_share_index ~ 1 + (1 ii) + (1 jj) + (1 kk))</code>
M1	<code>bbrm(z_post_share_index ~ 1 + z_kk_headline_believable + z_kk_headline_controversial + z_kk_headline_surprising + z_kk_headline_interesting + z_kk_headline_positive + z_kk_headline_true + z_headline_believable + z_headline_controversial + z_headline_surprising + z_headline_interesting + z_headline_positive + (1 ii) + (1 jj) + (1 + z_headline_believable + z_headline_controversial + z_headline_surprising + z_headline_interesting + z_headline_positive kk))</code>
M2	<code>brm(z_post_share_index ~ 1 + z_ii_post_controversial + z_ii_post_surprising + z_ii_post_interesting + z_ii_post_positive + z_post_controversial + z_post_surprising + z_post_interesting + z_post_positive + (1 + z_post_controversial + z_post_surprising + z_post_interesting + z_post_positive ii) + (1 jj) + (1 kk))</code>
M3	<code>brm(z_post_share_index ~ 1 + z_post_agree*z_post_align + z_ii_post_agree*z_ii_post_align + (1 + z_post_agree*z_post_align ii) + (1 jj) + (1 kk))</code>
M4	<code>brm(z_post_share_index ~ 1 + x_post_framing*z_jj_mfq_indi + x_post_framing*z_jj_mfq_bind + x_post_framing*z_jj_mfq_prop + (1 jj) + (1 + x_post_framing*z_jj_mfq_indi + x_post_framing*z_jj_mfq_bind + x_post_framing*z_jj_mfq_prop kk))</code>
M5	<code>brm(z_post_share_index ~ 1 + x_post_framing*z_jj_conservatism + (1 jj) + (1 + x_post_framing*z_jj_conservatism kk))</code>

Note. Table provides an overview of the R formulas of the models in this Study. “kk” indicates headline-level, “ii” indicates post-level, and “jj” indicates a user-level variable. See the R code for details on the implementation, such as number of chains, warm-up, etc.

Table S19*R Formulas for Bayesian Multilevel Linear Regression Models (Mediation)*

Model	R Formula
Mediation	<pre>brm(bf(z_post_agree ~ 1 + z_headline_true + x_post_framing*z_jj_mfq_indi + x_post_framing*z_jj_mfq_bind + (1 jj) + (1 + x_post_framing*z_jj_mfq_indi + x_post_framing*z_jj_mfq_bind kk)) + bf(z_post_align ~ 1 + z_headline_true + x_post_framing*z_jj_mfq_indi + x_post_framing*z_jj_mfq_bind + (1 jj) + (1 + x_post_framing*z_jj_mfq_indi + x_post_framing*z_jj_mfq_bind kk)) + bf(z_post_share_index ~ 1 + z_headline_true + x_post_framing*z_jj_mfq_indi + x_post_framing*z_jj_mfq_bind + x_post_framing*z_post_agree*z_post_align + (1 jj) + (1 + x_post_framing*z_jj_mfq_indi + x_post_framing*z_jj_mfq_bind + x_post_framing*z_post_agree*z_post_align kk)))</pre>

Note. Table provides an overview of the R formula of the mediation model in this Study. “kk” indicates headline-level, “ii” indicates post-level, and “jj” indicates a user-level variable. See the R code for details on the implementation, such as number of chains, warm-up, etc.

Study 2

Table S20

R Formulas for Bayesian Multilevel Linear Regression Models (Mediation)

Model	R Formula
Deliberation	<pre>brm(bf(z_post_deliberation_index ~ 1 + x_headline_true + z_headline_familiar + z_jj_crt*x_post_framing*z_jj_mfq_indi + z_jj_crt*x_post_framing*z_jj_mfq_bind + (1 jj) + (1 + z_jj_crt*x_post_framing*z_jj_mfq_indi + z_jj_crt*x_post_framing*z_jj_mfq_bind + z_headline_familiar kk)) + bf(z_post_share_index ~ 1 + z_headline_familiar + x_post_framing*z_jj_mfq_indi + x_post_framing*z_jj_mfq_bind + x_headline_true*x_post_framing*z_post_deliberation_index + (1 jj) + (1 + x_post_framing*z_jj_mfq_indi + x_post_framing*z_jj_mfq_bind + x_post_framing*z_post_deliberation_index + z_headline_familiar kk)))</pre>
Response Time	<pre>brm(bf(z_post_response_time ~ 1 + x_headline_true + z_headline_familiar + z_jj_crt*x_post_framing*z_jj_mfq_indi + z_jj_crt*x_post_framing*z_jj_mfq_bind + (1 jj) + (1 + z_jj_crt*x_post_framing*z_jj_mfq_indi + z_jj_crt*x_post_framing*z_jj_mfq_bind + z_headline_familiar kk)) + bf(z_post_share_index ~ 1 + z_headline_familiar + x_post_framing*z_jj_mfq_indi + x_post_framing*z_jj_mfq_bind + x_headline_true*x_post_framing*z_post_response_time + (1 jj) + (1 + x_post_framing*z_jj_mfq_indi + x_post_framing*z_jj_mfq_bind + x_post_framing*z_post_response_time + z_headline_familiar kk)))</pre>

Note. Table provides an overview of the R formula of the mediation models in this Study. “kk” indicates headline-level, “ii” indicates post-level, and “jj” indicates a user-level variable. See the R code for details on the implementation, such as number of chains, warm-up, etc.

5 Additional Analyses of Inferring Political Ideology via Stance on Vaccination

In Study 3 we determined users’ political ideology via their past retweeting behavior and the accounts they follow (i.e., based on what verified political accounts a user follows and shares). However, this approach might be limited because it requires to have knowledge about the ideology about individuals in a user’s network at a given time and mainly relies on a number of known political accounts (usually famous politicians, pundits, or organizations). Instead one could also determine a users ideology based on the contents they share specific to the conversation at hand, here content that is “anti-vax” or “pro-vax”. The benefit of this approach is that users’ ideology can be determined directly from their behavior in the course of relevant posts, such as posts about vaccinations, and during the same time frame that the data is being collected. Of course, this in turn relies on the topic being sufficiently politically polarized, i.e., “anti-vax” being strongly correlated with conservatism and “pro-vax” being strongly correlated with being liberal (Clarkson & Jasper, 2022; Jiang et al., 2021; Kerr et al., 2021; Stroope et al., 2021).

In a robustness check to the analysis to Study 3, we infer ideology via users’ stance on vaccination and test whether ideology measured via the misinformation-exposure API explains as much or more variance compared to our proxy. We inferred each user’s position on COVID vaccination and mandates. More specifically, we employed an unsupervised stance detection method (Darwish et al., 2020) which uses dimensionality reduction to project users onto a low-dimensional space, followed by clustering, that allows identifying representative core users. To classify the stance of each user in the corpus as either pro-vaccination (“pro-vax”) or anti-vaccination (“anti-vax”), we compute the cosine similarity between each pair of users based on (1) (re-)tweeting identical tweets; (2) the hashtags that users use; and (3) the accounts they retweet. We then refit the models from Study 3 using stance instead of political ideology and fit a model that adds stance to the ideology x moral values model, and compare whether ideology explains as much or more

variance compared to stance and whether stance has any explanatory power beyond ideology.

We find that a model using political ideology has significantly higher explanatory power than a model using stance ($\Delta_{ELPD} = 20.06, SE = 5.85, z = 3.43$) and stance did not add any explanatory power to the ideology x moral values model ($\Delta_{ELPD} = 22.31, SE = 14.08, z = 1.58$), indicating that stance was indeed simply a proxy for ideology. Furthermore, this robustness check adds additional credibility to the political ideology score determined by the misinformation-exposure API by showing alignment with alternative measures of ideology.

6 Additional Analyses of Analytical Thinking

Study 2 replicated the results of Study 1 and tested whether lack of deliberation could be an alternative explanation. We tested whether an alignment of moral values and framing distracted participants from post accuracy and plausibility (via reducing deliberation), leading to more sharing of misinformation. We did not find that deliberation of posts mediated the effect of aligning a participant’s moral values and a post’s framing. To further strengthen these findings, we directly tested whether deliberation reduced sharing of false (vs true) news by fitting four additional linear regression models.

First, we fitted model M7 that predicted sharing intentions as a function of analytical thinking, headline veracity and their interaction. This model tested whether analytical thinking (CRT-2) reduced sharing of false (vs true) news. We found no effect for this relationship ($\beta = 0.03, [-0.04, 0.08]$). Second, we fit model M8 that predicted sharing intentions as a function of deliberating over sharing a post, headline veracity and their interaction. This model tested whether deliberation, over each sharing behavior, reduced sharing of fake (vs true) news. We, again, found no evidence for this relationship ($\beta = 0.01, [-0.09, 0.11]$). Furthermore, these models did not predict sharing intentions more accurately than the null model

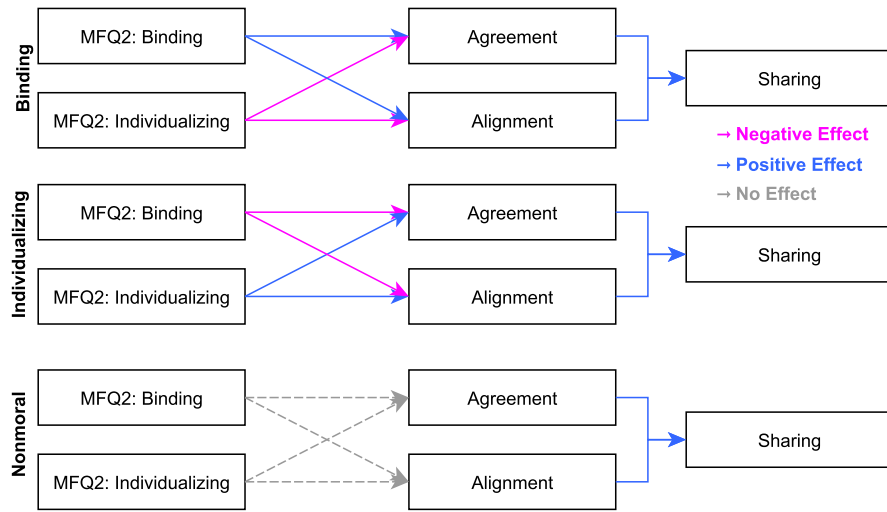
($\Delta_{ELPD} = -2.61$, $SE = 8.52$, $z = -0.31$; $\Delta_{ELPD} = 7.70$, $SE = 9.04$, $z = 0.85$), which predicted sharing intentions as a function of random intercepts for participants, headlines, and posts. Third, we fitted model M9 that predicted sharing intentions as a function of headline veracity, analytical thinking (CRT-2), headline believability and their interaction. This model tested whether analytical thinking increased accuracy and plausibility considerations, as argued by past work (Pennycook & Rand, 2019b). If analytical thinking results in participants estimating and considering plausibility, then there should be a positive interaction of CRT and headline believability, meaning that analytical thinkers should have higher plausibility concerns than “lazy” thinkers (Pennycook & Rand, 2019b). We found no effect for this relationship ($\beta = -0.03$, $[-0.08, 0.02]$). Fourth, focusing only on nonmoral stimuli, we fitted model M10 that predicted sharing intentions as a function of analytical thinking, headline veracity, and their interaction. This model investigated whether the failure to replicate past findings of analytical thinking reducing misinformation sharing was caused by most of our stimuli being moralized (two thirds). It could be that for moral-emotional stimuli accuracy concerns were superseded by participants’ intuitions of right and wrong. Supporting this idea, we find that, for nonmoral stimuli, analytical thinking reduces sharing of misinformation ($\beta = -0.10$, $[-0.17, -0.02]$) but not true information ($\beta = -0.04$, $[-0.14, 0.05]$).

We further replicated the mediation analysis from Study 1, which mediated the effect of matching a participant’s moral values and a post’s moral framing on sharing intentions via agreement and moral alignment with the post. Similar to Study 1, we compared across the three moral framing conditions the total indirect effects of participants’ endorsement of Binding and Individualizing values on sharing intentions via their ratings of how much they agreed with the post, how much the post aligned with their moral values, and the interaction of the two ratings, while controlling for headline veracity and, additionally, headline familiarity. Figure S5 provides an overview of the observed relationships. Supporting the original Hypothesis 2 in Study 1, we found that participants’

endorsement of Binding values had a positive indirect effect on sharing intentions in the Binding framing condition ($\beta = .21, [.15, .27]$) but a negative indirect effect in the Individualizing framing condition ($\beta = -.10, [-.15, -.05]$). Furthermore, participants' endorsement of Individualizing values had a positive indirect effect on sharing intentions in the Individualizing framing condition ($\beta = .21, [.15, .28]$) but a negative indirect effect in the Binding framing condition ($\beta = -.08, [-.14, -.03]$). Lastly, participants' endorsement of Binding ($\beta = .00, [-.05, .06]$) and Individualizing ($\beta = .03, [-.02, .09]$) values had no indirect effect in the nonmoral framing condition. These findings, again, support the idea of motivational drivers of message sharing.

Figure S5

Replication of the mediation in Study 1: Effect of matching moral values and framing via perceived agreement and alignment with the post



Note. The figure shows a clear separation between the effect of matching moral framing, which leads to an increase of sharing intentions (blue color), mismatched moral framing, which leads to a decrease of sharing intentions (red color), and not addressing moral values, which has no effect on sharing intentions (grey color), across all conditions.

7 Additional Analyses of Stimuli Order Effects

In our studies 1 & 2, we presented participants repeatedly with stimuli and items to rate said stimuli (e.g., believability, agreement, etc). Participants were then asked for their

sharing intentions. Thus, participants might have learned after the first trial that they will have to indicate their sharing intentions for the presented social media posts at each trial. In the following iterations participants might have decided about their sharing intentions during post presentation (before all ratings) and this could have influenced their subsequent stimuli ratings (i.e., rating to justify their sharing intentions; e.g., rate a headline as more believable if they want to share the post). Therefore, as a robustness check, we conducted an additional analysis to test for potential order effects, that is whether the effect of our ratings on sharing intentions increased over trial iterations. We ran the respective models (M1, M3, M6) again while controlling for stimulus order but found no significant interactions of stimulus order and main predictors (M1:

$\beta_{familiar:order} = -0.00, [-0.02, 0.01]$), M3: $\beta_{agreement:order} = 0.00, [-0.03, 0.02]$; M6: $\beta_{order:framing1:individualizing} = 0.00, [-0.02, 0.02]$, $\beta_{order:framing2:binding} = 0.01, [-0.02, 0.03]$).

8 Additional Analyses of Public vs Private Sharing

In our studies 1 & 2, we utilized an index that combined public and private sharing intentions (public sharing, public liking, private online sharing, private offline sharing). While this index had high reliability above 0.8, the underlying items might tap into distinct motivations of sharing. Given people’s motivation to express their values and other potential social functions of public sharing, such as aligning with group values and identities (Oh & Syn, 2015) it could be that moral alignment is primarily drives public sharing intentions and not private and offline sharing which do not fulfill the same sharing motivations (e.g., difficult to signal to the group through private or offline sharing).

Thus, we refit the main model $M4_{veracity}$ (which includes moral alignment and its interaction with veracity) from Study 1 and Study 2 respectively using each index as a separate outcome variable and analyze whether the observed effects of moral alignment drives either hold for public, private sharing, or both.

Study 1

We find that the previously reported effects for moral alignment only hold up for public sharing. In this case, model $M4_{\text{veracity}}$ predicted sharing intentions significantly more accurate compared to the baseline model M0 ($\Delta_{ELPD} = 66.64, SE = 17.32, z = 3.85$) and model M5 ($\Delta_{ELPD} = 33.67, SE = 17.18, z = 1.96$) which uses political ideology instead of moral values. Consistent with the results in Study 1, we find that participants Binding values predicted significantly more public sharing intentions for false posts framed with Binding compared to Individualizing values $\Delta\beta = 0.19, [0.10, 0.29]$ and, to a lesser extent, non-moral posts $\Delta\beta = 0.09, [-0.01, 0.18]$. For true information however, participants' endorsement of Binding values did not predict greater sharing intentions in the Binding framing condition than in the Individualizing framing condition ($\Delta\beta = 0.09, [-0.05, 0.23]$) or in the nonmoral framing condition ($\Delta\beta = 0.09, [-0.05, 0.22]$). In other words, participant showed higher sharing intentions for sharing misinformation (but not true information) framed with Binding values (aligned) than posts with Individualizing values (misaligned). Likewise, participants' endorsement of Individualizing values predicted greater sharing intentions, for false posts, in the Individualizing framing condition than in the Binding framing condition ($\Delta\beta = 0.20, [0.11, 0.28]$) and in the nonmoral framing condition ($\Delta\beta = 0.11, [0.01, 0.21]$). For true information, participants' endorsement of Individualizing values predicted greater sharing intentions in the Individualizing framing condition than in the Binding framing condition ($\Delta\beta = 0.18, [0.06, 0.30]$) and, to a lesser extent, in the nonmoral framing condition ($\Delta\beta = 0.09, [-0.05, 0.23]$). In other words, participants with Individualizing values had greater sharing intentions for misinformation framed with Individualizing values (aligned) than nonmoral posts (neutral) and posts with Binding values (misaligned) but for true information this effect was dampened and participants had only greater sharing intentions for posts framed with Individualizing values (aligned) vs Binding (misaligned). Note again that the effect sizes of moral alignment were, across all conditions, lower for true information compared to

misinformation even when the effects were still significant (e.g., Individualizing vs Binding framing for participants with Individualizing values), indicating a generally lower sensitivity of true information to moral alignment.

Table S21

Effect of participant values on public online sharing across framing conditions and stimuli veracity

	Aligned vs Misaligned		Aligned vs Neutral	
	False	True	False	True
Binding				
Values	0.19 [0.10, 0.29]	0.09 [-0.05, 0.23]	0.09 [-0.01, 0.18]	0.09 [-0.05, 0.22]
Individualizing				
Values	0.20 [0.11, 0.28]	0.18 [0.06, 0.30]	0.11 [0.01, 0.21]	0.09 [-0.05, 0.23]

Note. Table shows the difference in effect of moral alignment vs misalignment or non-moral (neutral) framing on public sharing intentions for both false and true posts. Table shows that the effect of moral alignment is generally larger for false posts than for true posts and that the effect of moral alignment is generally not significant for true posts, indicating that misinformation is more sensitive to moral alignment.

For private online and offline sharing, model $M4_{\text{veracity}}$ did not predict sharing intentions more accurately than the baseline model M0 ($z_{\text{private}} = 1.28$; $z_{\text{offline}} = -1.64$) or model M5 ($z_{\text{private}} = 1.42$; $z_{\text{offline}} = -2.12$) which uses political ideology instead of moral values. Furthermore, we did not find an effect of moral alignment on sharing of either true or false posts, except for aligned Individualizing values vs non-moral and misaligned posts. See Table S22 and Table S23 for an overview of the effects across all conditions and stimuli veracity for private online and offline sharing respectively.

Study 2

We find that the previously reported effects for moral alignment only hold up for public sharing. In this case, model $M4_{\text{veracity}}$ predicted sharing intentions significantly more accurate compared to the baseline model M0 ($\Delta_{\text{ELPD}} = 63.66$, $SE = 16.74$, $z = 3.80$) and model M5 ($\Delta_{\text{ELPD}} = 48.91$, $SE = 16.91$, $z = 2.89$) which uses political ideology instead of moral values. Consistent with the results in Study 1, we find that participants Binding

Table S22

Effect of participant values on private online sharing across framing conditions and stimuli veracity

	Aligned vs Misaligned		Aligned vs Neutral	
	False	True	False	True
Binding				
Values	0.02 [-0.07, 0.11]	0.03 [-0.11, 0.16]	0.00 [-0.10, 0.10]	0.03, [-0.11, 0.16]
Individualizing				
Values	0.11 [0.03, 0.18]	0.09 [-0.02, 0.20]	0.10 [0.02, 0.18]	0.07 [-0.04, 0.19]

Note. Table shows the difference in effect of moral alignment vs misalignment or non-moral (neutral) framing on public sharing intentions for both false and true posts. Table shows no effect of moral alignment except for Individualizing values on Individualizing vs Binding and non-moral framing for false posts. Results indicate that private sharing is not driven by moral alignment.

Table S23

Effect of participant values on private offline sharing across framing conditions and stimuli veracity

	Aligned vs Misaligned		Aligned vs Neutral	
	False	True	False	True
Binding				
Values	0.03 [-0.07, 0.12]	0.02 [-0.11, 0.15]	0.02 [-0.08, 0.11]	0.08 [-0.05, 0.21]
Individualizing				
Values	0.07 [-0.00, 0.15]	0.02 [-0.09, 0.13]	0.01 [-0.07, 0.09]	-0.04 [-0.16, 0.08]

Note. Table shows the difference in effect of moral alignment vs misalignment or non-moral (neutral) framing on public sharing intentions for both false and true posts. Table shows no effect of moral alignment except across all conditions and for both true and false posts. Results indicate that private offline sharing is not driven by moral alignment.

values predicted significantly more public sharing intentions for false posts framed with Binding compared to Individualizing values $\Delta\beta = 0.18, [0.09, 0.26]$) and non-moral posts $\Delta\beta = 0.12, [0.03, 0.21]$). For true information however, participants' endorsement of Binding values did not predict greater sharing intentions in the Binding framing condition than in the Individualizing framing condition ($\Delta\beta = 0.12[-0.00, 0.24]$) or in the nonmoral framing condition ($\Delta\beta = 0.11[-0.02, 0.24]$). In other words, participant showed higher sharing intentions for sharing misinformation (but not true information) framed with Binding values (aligned) than non-moral posts (neutral) and posts with Individualizing values (misaligned). Likewise, participants' endorsement of Individualizing values predicted greater sharing intentions, for false posts, in the Individualizing framing condition than in the Binding framing condition ($\Delta\beta = 0.19[0.11, 0.26]$) and in the nonmoral framing condition ($\Delta\beta = 0.14, [0.06, 0.22]$). For true information, participants' endorsement of Individualizing values predicted greater sharing intentions in the Individualizing framing condition than in the Binding framing condition ($\Delta\beta = 0.12, [0.01, 0.22]$) and, albeit not significantly, in the nonmoral framing condition ($\Delta\beta = 0.12[-0.00, 0.23]$). In other words, participants with Individualizing values had greater sharing intentions for misinformation framed with Individualizing values (aligned) than nonmoral posts (neutral) and posts with Binding values (misaligned) but for true information this effect was dampened and participants had only greater sharing intentions for posts framed with Individualizing values (aligned) vs Binding (misaligned). Note again that the effect sizes of moral alignment were, across all conditions, lower for true information compared to misinformation even when the effects were still significant (e.g., Individualizing vs Binding framing for participants with Individualizing values), indicating a generally lower sensitivity of true information to moral alignment.

For private online and offline sharing, model $M4_{\text{veracity}}$ did not predict sharing intentions more accurately than the baseline model M0 ($z_{\text{private}} = 0.35; z_{\text{offline}} = 0.25$) or model M5 ($z_{\text{private}} = 0.54; z_{\text{offline}} = 0.54$) which uses political ideology instead of moral

Table S24*Effect of participant values on public online sharing across framing conditions and stimuli veracity*

	Aligned vs Misaligned			Aligned vs Neutral	
	False	True		False	True
Binding					
Values	0.18 [0.09, 0.26]	0.12 [-0.00, 0.24]		0.12 [0.03, 0.21]	0.11 [-0.02, 0.24]
Individualizing					
Values	0.19 [0.11, 0.26]	0.12 [0.01, 0.22]		0.14 [0.06, 0.21]	0.12 [-0.00, 0.23]

Note. Table shows the difference in effect of moral alignment vs misalignment or non-moral (neutral) framing on public sharing intentions for both false and true posts. Table shows that the effect of moral alignment is generally larger for false posts than for true posts and that the effect of moral alignment is only significant for Individualizing values and framing.

values. In line with the results for Study 1, we found no effect of moral alignment on private offline sharing, supporting our previous conclusion about separate underlying mechanisms for this kind of sharing. Similar to Study 1, we find an effect of moral alignment (vs misalignment) for private online sharing of misinformation but not true information similar to the results of public sharing. However, the effect sizes are smaller compared to those for public sharing and there is no effect for moral alignment compared to non-moral (neutral) posts. These findings indicate that there might distinct underlying dynamics driving public online sharing, private online sharing, and private offline sharing, with moral alignment mainly facilitating public online sharing and partially facilitating private online sharing but not offline sharing. See Table S25 and Table S26 for an overview of the effects across all conditions and stimuli veracity for private online and offline sharing respectively.

Table S25

Effect of participant values on private online sharing across framing conditions and stimuli veracity

	Aligned vs Misaligned		Aligned vs Neutral	
	False	True	False	True
Binding				
Values	0.10 [0.01, 0.19]	0.06 [-0.07, 0.19]	0.06 [-0.03, 0.16]	0.05, [-0.08, 0.18]
Individualizing				
Values	0.14 [0.07, 0.22]	0.10 [-0.00, 0.20]	0.10 [0.02, 0.17]	0.08 [-0.03, 0.19]

Note. Table shows the difference in effect of moral alignment vs misalignment or non-moral (neutral) framing on public sharing intentions for both false and true posts. Table shows an effect of moral alignment vs misalignment for false but not true posts. Results indicate that moral alignment might facilitate private online sharing of misinformation.

Table S26

Effect of participant values on private offline sharing across framing conditions and stimuli veracity

	Aligned vs Misaligned		Aligned vs Neutral	
	False	True	False	True
Binding				
Values	0.06 [-0.03, 0.15]	0.01 [-0.12, 0.14]	0.04 [-0.05, 0.14]	0.05 [-0.08, 0.18]
Individualizing				
Values	0.04 [-0.03, 0.12]	0.01 [-0.10, 0.11]	0.07 [-0.00, 0.16]	0.07 [-0.03, 0.18]

Note. Table shows the difference in effect of moral alignment vs misalignment or non-moral (neutral) framing on public sharing intentions for both false and true posts. Table shows no effect of moral alignment except across all conditions and for both true and false posts. Results indicate that private offline sharing is not driven by moral alignment.

9 Additional Models

Study 1

Table S27

Effect sizes for Model 4 when controlling for headline veracity

Values - Condition	β [CI]	$\Delta\beta_{no\ control}$
Binding - Binding	.26 [.16, .36]	0.0
Binding - Individualizing	.14 [.04, .24]	0.0
Binding - Nonmoral	.20 [.10, .30]	0.0
Individualizing - Binding	.07 [-.01, .14]	0.0
Individualizing - Individualizing	.23 [.16, .26]	0.0
Individualizing - Nonmoral	.14 [.06, .21]	0.0
Proportionality - Binding	.00 [-.09, .09]	0.0
Proportionality - Individualizing	-.05 [-.14, .04]	0.0
Proportionality - Nonmoral	-.03 [-.12, .07]	0.0
$\Delta\beta_{Binding:Binding-Individualizing}$	.12 [.03, .21]	0.0
$\Delta\beta_{Binding:Binding-Nonmoral}$	.06 [-.04, .15]	0.0
$\Delta\beta_{Individualizing:Individualizing-Binding}$	.17 [.09, .24]	0.01
$\Delta\beta_{Individualizing:Individualizing-Nonmoral}$	.10 [.01, .18]	0.0

Note. Table shows the effect sizes for a given moral value in a given framing condition, as well as the difference in effect sizes for a moral value across framing conditions (last 4 rows). Table shows that the reported effect sizes in Study 1 hold when controlling for headline veracity.

Study 2

Table S28

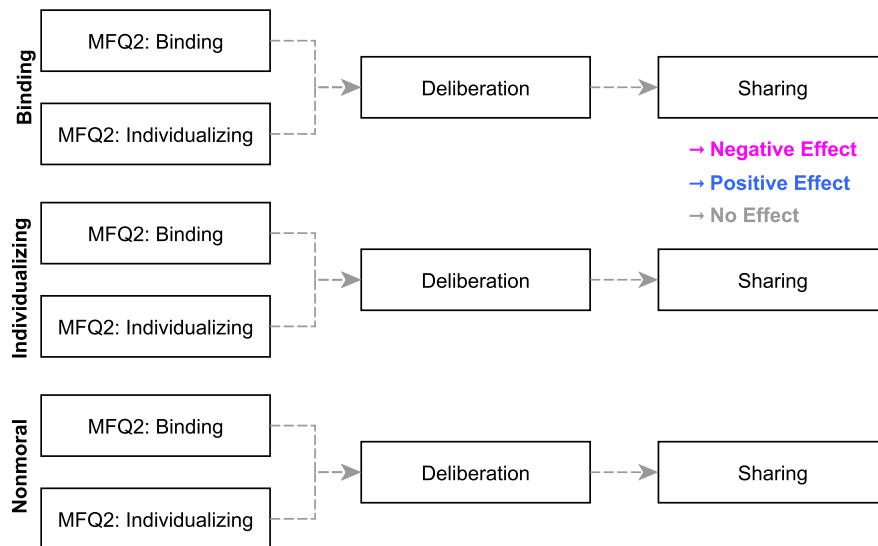
Effect sizes for Model 4 when controlling for headline veracity and familiarity

Values – Framing Condition	β [CI]	$\Delta\beta_{no\ control}$
Binding – Binding	.25 [.17, .33]	-0.01
Binding – Individualizing	.09 [.01, .18]	-0.02
Binding – Nonmoral	.14 [.05, .23]	0.01
Individualizing – Binding	.09 [.02, .16]	-0.02
Individualizing – Individualizing	.24 [.16, .32]	-0.02
Individualizing – Nonmoral	.12 [.05, .20]	-0.01
Proportionality – Binding	.01 [-.08, .10]	0.00
Proportionality – Individualizing	.00 [-.09, .09]	0.01
Proportionality – Nonmoral	.03 [-.06, .11]	0.01
$\Delta\beta_{Binding:Binding-Individualizing}$	.16 [.07, .24]	0.02
$\Delta\beta_{Binding:Binding-Nonmoral}$	.11 [.02, .19]	0.00
$\Delta\beta_{Individualizing:Individualizing-Binding}$	.15 [.08, .22]	0.00
$\Delta\beta_{Individualizing:Individualizing-Nonmoral}$	.12 [.04, .19]	-0.01

Note. Table shows the effect sizes for a given moral value in a given framing condition, as well as the difference in effect sizes for a moral value across framing conditions (last 4 rows). Table shows that the reported effect sizes in Study 2 hold when controlling for headline veracity.

Figure S6

Results from the preregistered mediation analysis of the effect of aligning moral framing and moral values via deliberation



Note. Figure shows that deliberation does not mediate the effect of moral alignment on sharing intentions. Importantly, there is no effect of deliberation on sharing intentions.