

Manuscript accepted for publication in
American Psychologist

**Justifying Antisemitism:
Political Liberalism and Perceptions of Prejudices**

Jordan W. Moon ^{1,2*}, Michael Barlev ³, & Steven L. Neuberg ³

¹ School of Psychology, University of Southampton

² Centre for Culture and Evolution, Brunel University of London

³ Department of Psychology, Arizona State University

* **Correspondence:** Jordan W. Moon, School of Psychology, University of Southampton, B44 University Road, Southampton, United Kingdom, SO17 1BJ.

Email: j.w.moon@soton.ac.uk

Author contributions (CRediT): Conceptualization (all authors), formal analysis (JWM), investigation (JWM, MB), methodology (all authors), visualization (JWM), writing—original draft (JWM), writing—review and editing (all authors).

Conflict of Interest Statement: The authors declare no conflicts of interest with respect to the authorship or publication of this article.

Financial Support: The studies reported here were supported by funding to Steven L. Neuberg from the Arizona State University Foundation for a New American University.

Abstract

Antisemitism has surged in the U.S. since the onset of the Israel– Hamas War. Yet, responses from liberals—who generally vigorously defend marginalized groups against prejudice and discrimination—have been tepid, leading some to suggest that liberals tolerate antisemitism. In three pre-registered experiments ($N = 979$), we investigate how Americans—liberals and conservatives—perceive antisemitism, and whether their perceptions depend on how the antisemitism is justified. We find that, absent justifications, individuals expressing antisemitism (and other prejudices) are generally disliked, and more so by liberals than by conservatives. However, when these individuals justify their antisemitism by disapproval of Israel and the war in Gaza, or violations by Israel of the human rights of Palestinians, they are liked more by liberals (but not conservatives). We find support for two group-based explanations for this “licensing” effect of justifications: Liberals evaluate individuals who express antisemitism and other prejudices more positively to the extent that the justifications they express (1) identify them as liberals (ingroup favoritism) and (2) suggest that they are not generally bigoted and are therefore less of a threat to their political alliances (alliance politics). A fuller understanding of why antisemitism is presently not more broadly condemned requires considering how it is expressed, the social information contained in the expression, and the implications of this information for perceivers.

Keywords: Prejudice, Moral Licensing, Justification, Social Perception, Political Ideology

Public Significance Statement

Do liberals tolerate antisemitism? It depends. When expressions of antisemitism are not justified, liberals strongly dislike them (even more than conservatives do). However, some justifications make liberals tolerate expressions of antisemitism. We find support for two group-based explanations for this effect of justifications: ingroup favoritism and alliance politics. To combat antisemitism requires considering not only the psychology of individuals expressing antisemitism but also of social perceivers.

Justifying Antisemitism:

Political Liberalism and Perceptions of Prejudices

Since the beginning of the Israel– Hamas War in October 2023, antisemitic hate crimes have skyrocketed (FBI, 2024). The year preceding the war, Jews were the targets of 11% of all hate crimes (despite representing about 2% of the U.S. population); in the year following the war this increased to a staggering 27%. Although anti-Jewish hate crimes were already the highest per capita, they now surpass anti-Black hate crimes by raw numbers— even though there are more than six times as many Black people in America as compared to Jews. By comparison, hate crimes toward Muslims only increased from 2% to 3% in the same period.

Surprisingly, such soaring antisemitism has led to little outcry from American liberals, a group that has historically been greatly concerned with advancing civil rights and defending minorities against prejudices and discrimination. The current state of affairs is well illustrated by considering American universities. Since the beginning of the war, American universities have become both the flashpoint for protests and encampments— some in opposition to, and some in support of, Israel— and the focus of the national debate about antisemitism. Universities— which have been centers of liberalism since at least the 1960s and 70s (Gross, 2013)— have long been considered havens for Jewish students (Krasner, 2024). Yet, more than 10% of antisemitic hate crimes documented by the Anti-Defamation League (ADL) in the year following the war occurred on university campuses— a six-fold increase over the preceding year (ADL, 2024). A majority of Jewish students now report being highly concerned about antisemitism on their university campuses— not from the political right, but from the political left (Wright et al., 2023).

How have universities responded? In the face of hostile environments for Jewish students, many administrators have shown indifference or even contempt. For example, at one university, administrators allowed barricaded encampments to operate largely unchecked and made generic statements about “free speech” even as Jewish students were barred from walking freely across campus (The Associated Press, 2024). As one university

professor put it, academics who commonly denounce micro-aggression and dog whistles were ignoring the macro-aggressions and fog horns of antisemitism (Pettit, 2025).

The soaring antisemitism and seeming indifference by American liberals has led the Editorial Board of the New York Times (2025) to conclude that “Progressives reject many other forms of hate even as some tolerate antisemitism.” However, reality may be more nuanced than this. First, people rarely express prejudices without accompanying justifications—and some justifications can make prejudices seem more socially acceptable (Billig, 1988; Crandall & Eshleman, 2003; Effron & Knowles, 2015). Second, the justifications given on the political left (e.g., disapproval of Israel) differ from those traditionally given on the right (e.g., Jewish power and influence over financial markets, governments, and the media)—and both differ from historical justifications that are less commonly used in the US today (e.g., Jews crucified Christ, Jews are racially inferior) (Feldman, 2024; Lewis, 2005).

Overview

In three experiments, we explore how American liberals (and conservatives) evaluate different expressions of anti-Jewish prejudice (antisemitism). We follow a traditional definition of prejudice as a negative affective evaluation of a social group or of individual members of that group significantly based on their group membership (e.g., Crandall & Eshleman, 2003), regardless of its causes. We consider antisemitism absent justifications, as well as politically left-coded, politically right-coded, and historical justifications. Experiments 1 and 2 also evaluate, for purposes of comparison, the effects of justifications on expressions of prejudices toward other minority groups (e.g., Black people, Muslims).

After describing the nuances of how people perceive antisemitism, with a focus on antisemitism justified by disapproval of Israel and violations by Israel of the human rights of Palestinians, we evaluate two group-based explanations for the liberal tolerance of (some expressions of) antisemitism: (1) Ingroup favoritism and (2) alliance politics.

Explanation 1: Ingroup Favoritism

In the U.S., views of the Israel-Palestine conflict used to be bipartisan, with majorities of both Democrats and Republicans sympathizing with Israel (Brenan, 2025; Rynhold, 2020). By 2025, however, the conflict had become highly politicized, with 75% of Republicans but only 21% of Democrats reporting that their sympathies lie with the Israelis (only 10% of Republicans but 59% of Democrats report that their sympathies lie with the Palestinians; the rest say “both” or “neither”, or have no opinion) (Brenan, 2025).

Indeed, young liberals have been especially passionate protestors in support of Palestinians in Gaza and against Israel following the outbreak of the Israel– Hamas War (Silver, 2024)—in what has perhaps been the largest U.S. protest movement in history based on a foreign event (Chenoweth et al., 2024). Moreover, support for Palestine and sympathy for Palestinians have become intertwined with other liberal human rights causes, such as Black Lives Matter (Rynhold, 2020), and liberal organizations dedicated to other human rights causes (e.g., LGTB rights) have aligned with the pro-Palestinian movement (e.g., “Queers for Palestine”) (Chenoweth et al., 2024).

Given this, solidarity with Palestinians and opposition to Israel may currently serve as badges of being a political liberal. Consequently, liberals might tolerate individuals who express antisemitism *accompanied by such justifications* because they view them as being members of the same political coalition (Pietraszewski et al., 2015).

Explanation 2: Alliance Politics

Americans view prejudiced individuals harshly (Everett et al., 2019; Mae & Carlston, 2005; Sommers & Norton, 2006)—at least those prejudiced toward groups thought of as “marginalized” (Bergh & Brandt, 2022; Crandall et al., 2002). In fact, even people who agree with some prejudiced remarks evaluate those expressing them negatively (Mae & Carlston, 2005). Although people dislike expressions of prejudice for many reasons, one might be that they infer that prejudice toward one group is indicative of prejudices toward other groups as well—that is, of being generally bigoted (Everett et al., 2019).

But why are people so averse to bigotry in the first place? In modern societies, many large groups—political parties in particular—are alliances between multiple subgroups with

some—but far from all—overlapping interests. For example, at present, American liberals include Hollywood movie stars, urban professionals, Muslims, and LGBTQ minorities, whereas conservatives include wealthy business owners, evangelical Christians, police officers, and blue-collar White people. Although these alliances fail to cohere around a single unifying philosophy, they include *ad hoc* beliefs that function to propagandistically support the groups united under their banner (Pinsof et al., 2023). The ability of such alliances to hold together is greatly threatened by individual members of the alliance who might be prejudiced *indiscriminately* or *without cause*—and thus potentially against members of other subgroups within the alliance.

Perhaps because prejudiced individuals are so widely disliked, people are reluctant to express prejudices that they harbor unless they can provide socially acceptable justifications for them (Crandall & Eshleman, 2003). One aim of justifying a specific prejudice might be to mitigate the inference that one is generally bigoted. When people say that they are “not racist, but...”, they might be attempting to communicate that their negativity toward one group should not be interpreted as reflecting a broad pattern of intolerance toward all groups (Billig, 1988). That is, whereas expressing prejudice absent a justification runs the risk of implying a general pattern of intolerance toward different groups (Everett et al., 2019), successfully justifying this prejudice might reduce this inference.

Thus, from this perspective, an effective justification will reduce the perception that the individual expressing prejudice is bigoted. Justifications may not only reflect the true causes of prejudices, but also function as “*releasers*” that allow individuals to express prejudices they hold (Crandall & Eshleman, 2003; Hegarty & Golden, 2008; Norton et al., 2004; White II & Crandall, 2017). Here, we consider the possibility that the effects of specific justifications on liking have to do, in part, with the extent to which such individuals are thought to be bigoted.

Transparency and Openness

All experiments were pre-registered (Experiments 1-3: <https://osf.io/2yqsv/>; <https://osf.io/pfsqa/>; <https://osf.io/pfas2/>). All data, materials, and analysis code are

available on the Open Science Framework (<https://osf.io/mutxh/>). In response to reviewer feedback, reported analyses differ from the pre-registrations. All pre-registered analyses are reported in full in the Supplemental Materials. These experiments were determined to be exempt (low risk) by the Arizona State University IRB (STUDY00021025).

Method

Participants

In each of the three experiments reported here we recruited quota-based U.S. samples using Prolific (www.prolific.com), intended to mirror the US population in age, gender, ethnicity, and political affiliation (referred to as “representative sampling” by Prolific, though it likely does not achieve perfect demographic balance). The data were collected between October and December of 2024. As pre-registered, we excluded participants who failed checks for attention or AI use (10% or fewer in each sample). See Supplemental Materials for details about exclusion criteria.

After exclusions, Experiment 1 included 270 participants (129 men, 138 women) with an average age of 45 ($SD = 16$). Most participants were White (152), Black (50), or Hispanic/Latino (23). The sample included similar numbers of Democrats (93), Republicans (86), and Independents (89). On a liberalism-conservatism scale (0 = *Extremely liberal*, 100 = *Extremely conservative*), the average score was 49 ($SD = 32$). The most common religious affiliations were Christian (132), nonreligious (82), and “spiritual but not religious” (22). The numbers of Jewish participants in all three experiments were negligible. See Supplemental Materials for additional demographic information.

Experiment 2 included 277 participants after exclusions (139 men and 135 women), with an average age of 46 ($SD = 16$). Most participants were White (154), Black (44), or Hispanic/Latino (29). The sample included similar numbers of Democrats (98), Republicans (88), and Independents (89). On a liberalism-conservatism scale, the average score was 48 ($SD = 31$). The most common religious affiliations were Christian (138), nonreligious (68), and “spiritual but not religious” (34).

Experiment 3 included 432 participants after exclusions (223 men and 207 women) with an average age of 45 ($SD = 16$). Most participants were White (276), Black (68), or Hispanic/Latino (31). Experiment 3 asked about political party affiliation in a more fine-grained way, revealing that 34 participants identified as MAGA Republicans, 101 as Traditional Republican/Conservative, 25 as Moderate/Centrist, 100 as Traditional Democrat/Liberal, 32 as Progressive, and 134 as Independent. On a liberalism-conservatism scale, the average score was 49 ($SD = 32$). The most common religious affiliations were Christian (253), nonreligious (89), and “spiritual but not religious” (43).

Procedure

All three experiments followed the same basic procedure. After giving informed consent, participants read about one or more individuals (with randomly selected male or female names) expressing a prejudice toward a specific minority group, either with a justification or without. After evaluating the prejudiced individual(s) on three questions, participants answered several individual-difference items and provided demographic information.

Stimuli

Experiment 1 was a 3 (Target, within-subjects) $\times$ 2 (Justification, between-subjects) design. Participants read about three individuals prejudiced against Jews, Black people, and Muslims. To mask the study purpose, participants saw these in random order interspersed with five distractor profiles (e.g., abortion rights supporter, recreational drugs user). For each of the three profiles, participants were randomly assigned to see either a non-justified expression of prejudice (“[Name] doesn’t like [Jews/Black people/Muslims]”) or a justified one. The justifications were written to be ecologically valid—that is, to map onto justifications which are presently commonly given for anti-Jewish, anti-Black, and anti-Muslim prejudices.

- “[Name] doesn't like Jews, because [Name] strongly disapproves of Israel and its war in Gaza.”

- “[Name] doesn't like Black people, because [Name] believes Black people are more likely to commit crimes.”
- “[Name] doesn't like Muslims, because [Name] believes their culture and values conflict with Western culture and values.”

Experiment 2 was a 5 (Target, within-subjects) $\times$ 3 (Justification, between-subjects) design. Participants read about individuals prejudiced against Jewish, Black, Muslim, Chinese, and Russian Americans (in randomized order), again interspersed with five distractor profiles. Experiment 2 differed from Experiment 1 in two important ways. First, to make clear that the group toward whom the prejudices are expressed were from the United States, we added to all five groups the qualifier “Americans”.

Second, we balanced whether the justifications given were politically left- or right-coded, thereby having tighter experimental control between prejudices than in Experiment 1. For each of the five prejudiced individuals, participants were randomly assigned to see either a non-justified prejudice or a prejudice justified by (1) human rights violations by the group's country of origin or ancestral home (liberal-coded), or (2) threats to American traditional values (conservative-coded):

- “[Name] doesn't like Jewish Americans ... (1) because [Name] believes that Israel violates the human rights of Palestinians (2) because [Name] believes that they use their power and influence to undermine American traditional values.”
- “[Name] doesn't like Black Americans ... (1) because [Name] believes that many African countries violate the human rights of gay men and lesbians (2) because [Name] believes that they disregard and even drag down American traditional values.”
- “[Name] doesn't like Muslim Americans ... (1) because [Name] believes that many Muslim countries violate the human rights of women (2) because [Name] believes that they want to replace American traditional values with Islamic law.”

- “[Name] doesn’t like Chinese Americans (1) because [Name] believes that China violates the human rights of Muslim minorities (e.g., the Uyghurs) (2) because [Name] believes that they segregate themselves in communities that are indifferent to American traditional values.”
- “[Name] doesn’t like Russian Americans (1) because [Name] believes that Russia violates the human rights of Ukrainians (2) because [Name] believes that they covertly seek to undermine American traditional values.”

Finally, Experiment 3 was a fully between-subjects design with participants only reading about one individual prejudiced against Jews. Participants saw one of the following five conditions, each providing a different justification for disliking Jews: No justification, disapproval of Israel (“...strongly disapproves of Israel and its war in Gaza”), conspiracy (“...believes that Jews hold too much power and influence over financial markets, governments, and the media”), eugenics-race (“...believes that Jews are biologically distinct and genetically inferior to other races”), or religion (“...believes that Judaism is a heretical religion and that Jews carry responsibility for the crucifixion of Jesus Christ”).

Focal Dependent Measures

In all experiments, the prejudiced individuals were evaluated on the following three questions (in set order): (1) **Liking**: “To what extent do you like or dislike [Name]?” (1=*dislike extremely*, to 7=*like extremely*); (2) **Inferred liberalism-conservatism**: “How politically conservative or liberal do you think [Name] is?” (1=*very liberal*, to 7=*very conservative*); (3) **Inferred bigotry**: “How likely do you think it is that [Name] is a bigot (i.e., intolerant and prejudiced toward members of groups different from him/her)?” (1=*extremely unlikely*, to 7=*extremely likely*).

Experiment 3 asked three additional questions about the prejudiced individual (in set order): “How loyal versus disloyal is [Name] to members of his/her own group?” (1=*extremely disloyal*, 7=*extremely loyal*); “How kind versus selfish is [Name] to members of his/her own group?” (1=*extremely selfish*, 7=*extremely kind*); and “How kind versus selfish is [Name] to members of groups different from him/her?” 1= *extremely selfish*,

7=*extremely kind*). We report results on these additional questions in the Supplementary Materials.

Individual Differences

After evaluating the prejudiced individuals, participants answered several items: (1) feeling thermometers toward the prejudiced-against minority groups, from 0 (*extremely COLD and NEGATIVE*) to 100 (*extremely WARM AND POSITIVE*); (2) their own political liberalism-conservatism (0=*extremely LIBERAL*, 100=*extremely CONSERVATIVE*); and (3) their own views of the 2023 Israel– Hamas War: “Who is more responsible for the current conflict in Israel and Gaza?” (1=*Hamas is solely responsible*, 7=*Israel is solely responsible*).

Results

We first present the basic findings from all three experiments, focusing on the effects of justifications on liking, and how these effects differ for liberal and conservative perceivers.

All analyses in this manuscript were conducted using R, version 4.5.0 (R Core Team, 2023). Given the within-subjects design in Experiments 1 and 2 (i.e., participants saw multiple prejudiced profiles), we fitted linear mixed models (LMMs) using the *lme4* and *lmerTest* packages (Bates et al., 2015; Kuznetsova et al., 2017). We fitted the full Target (i.e., individuals expressing prejudice toward a specific group) $\times$ Justification $\times$ Participant Liberalism-Conservatism models with random intercepts by participant. Liberalism-conservatism was grand-mean centered prior to analysis. Models were estimated using restricted maximum likelihood (REML). In Experiment 3, participants only saw a single prejudiced individual (fully between-subjects), so we used OLS regression to estimate the effects of justification on liking, including the Justification $\times$ Participant Liberalism-Conservatism interaction.

In the main text, we focus on effects relevant to our hypotheses. Full model results (including all main effects and both 2- and 3-way interactions) can be found in Tables S1-S3 in the Supplemental Materials.

In all experiments, we use simple slopes analyses (Aiken & West, 1991) to test the effects of justifications on liking for liberal and conservative participants, defined as the

predicted effects among participants one standard deviation below and above the mean on the 0-100 liberalism-conservatism scale. Although it is possible to probe interactions at any level of a moderator, Aiken and West (1991) recommend estimating scores at one standard deviation above and below the mean, as there will nearly always be sufficient data at these values (i.e., the estimates are not extrapolations), and to maintain consistent guidelines across studies. As the present experiments used quota-based samples, these points on the liberalism-conservatism scale can reasonably be thought to represent respondents who are somewhat liberal and somewhat conservative relative to the US population.

In addition to simple slopes, we report Johnson-Neyman ranges of statistical significance. The Johnson-Neyman technique estimates the ranges of liberalism-conservatism for which the effect of each justification was statistically significant. Because this technique is not compatible with the multilevel models used here, it was computed using separate regression models for each justification, using the *interactions* package (Long, 2019) in R.

Experiment 1

Experiment 1 aimed to examine how justifications influence perceptions of individuals expressing antisemitism, by liberals and conservatives. Additionally, it examined the effects of justifications on perceptions of individuals expressing prejudices against two other marginalized groups: Black people and Muslims. In the U.S., like Black people, Jews are often viewed as a racial or ethnic minority, and like Muslims, Jews are a minority religion.

We note that the justifications used in this experiment were not written to be comparable between prejudices. Therefore, although we evaluate the effects of justifications on each of the three prejudices, we are limited in the inferences we can draw from comparing them. We return to this point in Experiment 2.

We fitted a 3 (Target, within-subjects) $\times$ 2 (Justification, between-subjects) $\times$ Participant Political Orientation (centered) LMM predicting liking. See Table S1 for full model results, including all interactions and main effects.

How do liberals versus conservatives evaluate prejudices absent justifications?

Absent justifications, liberals liked all three prejudiced targets less than did conservatives (Anti-Jewish: $b = 0.49$, $t[784] = 4.10$, $p < .001$; Anti-Black: $b = 0.57$, $t[772] = 5.03$, $p < .001$; Anti-Muslim: $b = 0.68$, $t[752] = 6.33$, $p < .001$). See Figure 1.

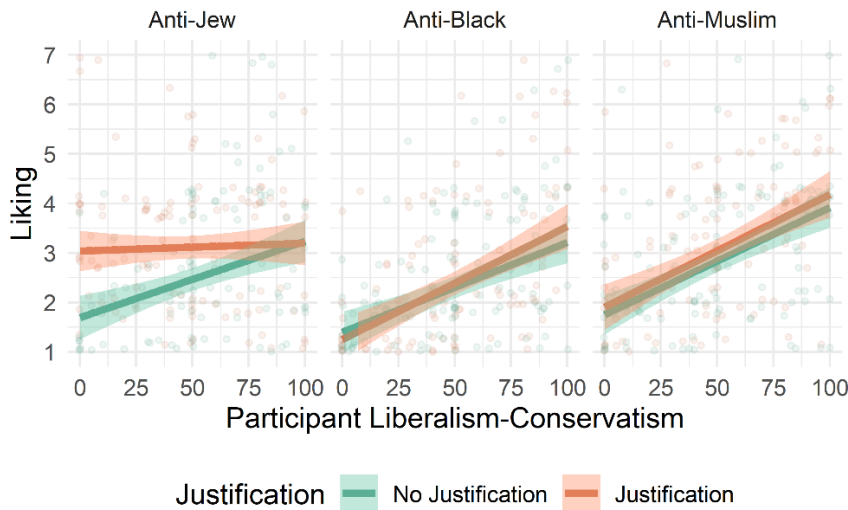


Figure 1. Experiment 1 liking ratings (“To what extent do you like or dislike [Name]?” from 1=dislike extremely, to 7=like extremely) as a function of the prejudice expressed, whether a justification was given, and participant liberalism-conservatism (0=extremely LIBERAL, to 100=extremely CONSERVATIVE). Ribbons represent 95% confidence intervals. Points are jittered to reduce overlap.

Do justifications “license” prejudices for liberals?

For the anti-Jewish targets, there was a significant Justification $\times$ Participant Liberalism-Conservatism interaction, $b = -0.43$, $t(681) = -2.85$, $p = .005$.

Unpacking this interaction with simple slopes analyses showed that when the anti-Jewish targets justified their prejudice by disapproval of Israel and the war in Gaza, liberals (modelled at 18 on the 0-100 liberalism-conservatism scale) liked them more, by 1.1 points on the 1-7 scale, $b = 1.10$, $t(681) = 5.04$, $p < .001$; however, this justification had no effect on how conservatives (modelled at 81 on the same scale) felt about the anti-Jewish targets, $b = 0.22$, $t(682) = 1.03$, $p = .305$. A Johnson-Neyman test (computed using an OLS regression,

with only the Anti-Jewish responses) suggested that this justification had a statistically significant effect ($p < .05$) for participants below 67 on the 0-100 liberalism-conservatism scale.

For the anti-Black and Muslim targets, the Justification $\times$ Participant Liberalism-Conservatism interactions were not statistically significant ($p = .334$ and $p = .827$, respectively).

Experiment 2

Experiment 1 found that the justification given for the anti-Jewish profile had different effects from the justifications given for the anti-Black and anti-Muslim profiles. However, the three justifications were conceptually very different and cannot directly be compared. Experiment 2 aimed to balance the justifications given, using both politically left- and right-coded justifications (human rights violations and threats to American traditional values, respectively). Experiment 2 assessed how justifications affect evaluations of individuals prejudiced against Jewish, Black, and Muslim Americans, as well as Chinese and Russian Americans. Chinese Americans were included because, like Jewish Americans, they are stereotyped as high in competence and low in warmth (Fiske et al., 2002). Russian Americans were included because Russia, like Israel, was at the time engaged in a highly polarizing war.

We fitted a 5 (Target: within-subjects) $\times$ 3 (Justification: between-subjects) $\times$ Participant Political Orientation (centered) LMM predicting liking. See Table S2 for full model results.

How do liberals versus conservatives evaluate prejudices absent justifications?

Replicating Experiment 1, absent justifications, liberals liked all five prejudiced targets less than did conservatives (Anti-Jewish: $b = 0.46$, $t[1317] = 3.54$, $p < .001$; Anti-Black: $b = 0.41$, $t[1331] = 3.04$, $p < .001$; Anti-Muslim: $b = 0.57$, $t[1321] = 4.34$, $p < .001$; Anti-Chinese: $b = 0.29$, $t[1272] = 2.36$, $p = .018$; Anti-Russian: $b = 0.32$, $t[1314] = 2.44$, $p = .015$). See Figure 2.

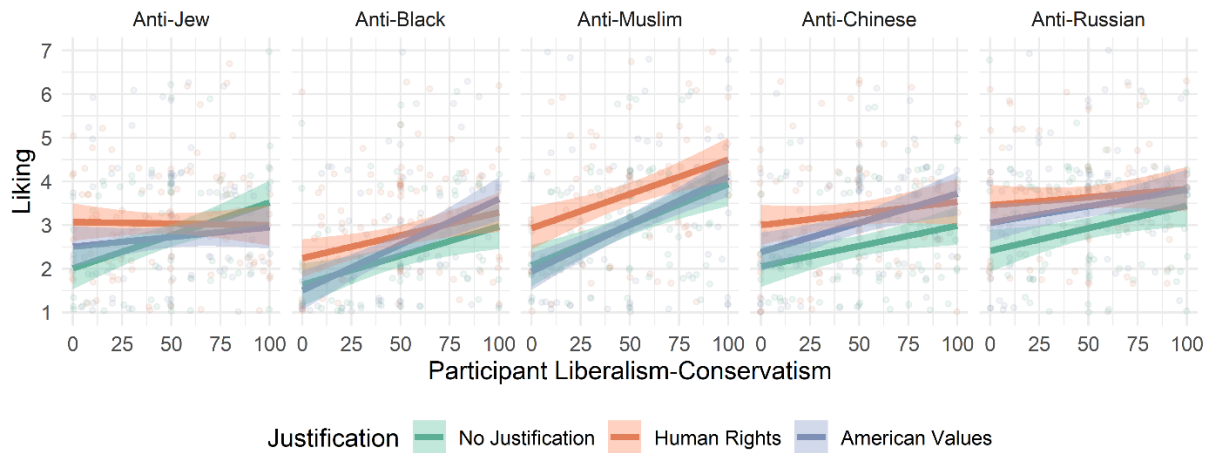


Figure 2. Experiment 2 liking ratings (“To what extent do you like or dislike [Name]?” from 1=dislike extremely, to 7=like extremely) as a function of the prejudice expressed, the justification given (including no justification), and participant liberalism-conservatism. Ribbons represent 95% confidence intervals. Points are jittered to reduce overlap.

Do left-coded (human rights) justifications “license” prejudices for liberals?

Yes. Replicating Experiment 1, for the anti-Jewish targets, there was a significant Justification (Human Rights vs. No Justification) $\times$ Participant Liberalism-Conservatism interaction, $b = -0.49$, $t(1169) = -3.03$, $p = .002$. When individuals justified their antisemitism via violations by Israel of the human rights of Palestinians, liberals (modelled at 17 on the 0-100 liberalism-conservatism scale) liked them more (by 0.78 points on the 1-7 scale), $b = 0.78$, $t(1168) = 3.44$, $p = .002$. However, this justification had no significant effect on how much conservatives (modelled at 79 on the same scale) liked the anti-Jewish targets, $b = 0.19$, $t(1169) = 0.82$, $p = .690$.

The left-coded human rights justifications had similar effects for other prejudiced targets. Although here the Justification $\times$ Participant Liberalism-Conservatism interactions were not statistically significant ($ps > .171$), estimating the effects of these justifications for liberals shows that they increased liking of targets with prejudices toward all targets: Black Americans: $b = 0.56$, $t(1168) = 2.41$, $p = .042$; Muslim Americans, $b = 0.80$, $t(1169) = 3.33$, $p = .003$, Chinese Americans, $b = 0.88$, $t(1168) = 3.76$, $p < .001$, and Russian Americans, $b = 0.93$, $t(1169) = 3.89$, $p < .001$.

We additionally evaluated whether left-coded justifications licensed prejudices for conservatives. Our findings here were less consistent. The human rights justifications increased liking of individuals prejudiced against Muslim Americans, $b = 0.62$, $t(1170) = 2.63$, $p = .023$ and Chinese Americans, $b = 0.63$, $t(1168) = 2.74$, $p = .017$, but not Black Americans, $b = -0.38$, $t(1169) = -1.68$, $p = .214$, or Russian Americans, $b = -0.52$, $t(1171) = -2.21$, $p = .071$.

Johnson-Neyman tests showed that the human rights justifications were statistically significant in the following liberalism-conservatism ranges: Between 9 and 31 for the anti-Jewish profile; between 3 and 73 for the anti-Black profile; at 20 and above for the anti-Muslim profile; at every value of liberalism-conservatism for the anti-Chinese profile; and between 3 and 82 for the anti-Russian profile.

Does every justification license prejudices for liberals?

No. The right-coded American traditional values justifications were markedly less effective in licensing prejudice for liberals—and conservatives. Although for the anti-Jewish targets there was a marginally significant Justification $\times$ Participant Liberalism-Conservatism (American Values vs. No Justification) interaction, $b = -0.33$, $t(1169) = -1.95$, $p = .051$, simple slopes analysis did not find that this justification made the anti-Jewish target liked more by either liberals, $b = 0.31$, $t(1169) = 1.30$, $p = .393$, or conservatives, $b = -0.34$, $t(1168) = -1.48$, $p = .303$.

For the four other prejudiced targets, the American traditional values justifications mostly had no effects on liking. There were no significant Justification $\times$ Participant Liberalism-Conservatism (American Values vs. No Justification) interactions ($ps > .171$), and simple slopes analyses found only two effects: for liberals, the American values justification increased liking for individuals prejudiced against Russian Americans, $b = 0.59$, $t(1168) = 2.62$, $p = .024$ (all other $ps > .189$), and for conservatives, for individuals prejudiced against Chinese Americans, $b = 0.65$, $t(1169) = 2.93$, $p = .010$ (all other $ps > .125$).

Johnson-Neyman tests found that the American values justification was statistically significant among participants above 20 for the anti-Chinese profile, and among participants

between 35 and 83 for the anti-Russian profile. It was not significant at any range for the anti-Jewish, anti-Black, or anti-Muslim profiles.

In sum, replicating Experiment 1, Experiment 2 revealed that, absent justifications, liberals tended to dislike individuals expressing antisemitism (and liked them less than conservatives did); however, if those individuals justified their prejudice via violations of the human rights of Palestinians, liberals—but not conservatives—liked them more.

Do similar left-coded justifications “license” prejudices toward other groups? Yes—for liberals. By matching justifications across prejudiced individuals, this experiment showed that for liberals (but mostly not for conservatives), justifying prejudices by appeals to human rights violations licensed prejudices broadly—toward Black, Muslim, Chinese, and Russian Americans. Was every justification effective in “licensing” antisemitism? No. A common right-coded justification—threats to American traditional values—did not lead liberals (or conservatives) to view individuals expressing antisemitism more favorably.

Experiment 3

Experiment 3 aimed to evaluate a larger number of justifications for antisemitism. As in Experiments 1 and 2, it compared antisemitism absent a justification to antisemitism justified by disapproval of Israel and the war in Gaza (identical wording to Experiment 1). We also included the following justifications: Jewish power and influence over financial markets, governments, and the media (a “conspiracy” justification, stronger wording than the Experiment 2 right-coded justification), and two historically common but no longer widespread justifications: eugenics-race and religion (cf. Feldman, 2024; Lewis, 2005).

We fitted a 5 (Justification: between-subjects) $\times$ Participant Political Orientation (centered) regression predicting liking. See Table S3 for full model results.

How do liberals versus conservatives evaluate antisemitism absent a justification?

As in Experiments 1 and 2, absent a justification, liberals liked the anti-Jewish target less than did conservatives, $b = 0.16$, $t(422) = 2.63$, $p = .009$. See Figure 3.

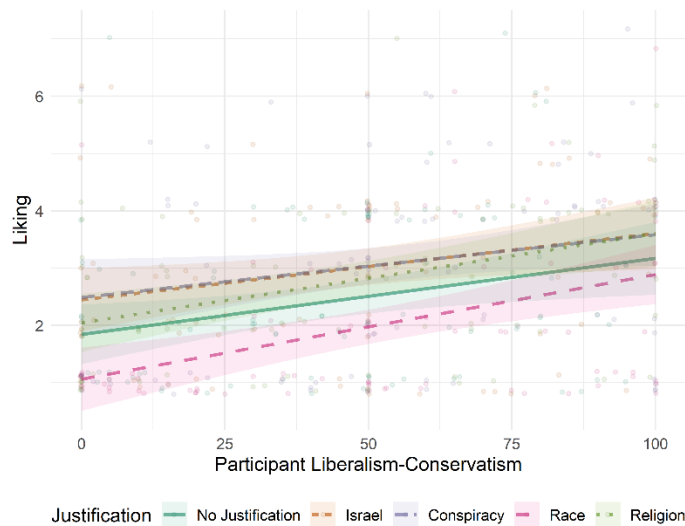


Figure 3. Experiment 3 liking ratings (“To what extent do you like or dislike [Name]?” from 1=*dislike extremely*, to 7=*like extremely*) of anti-Jewish individuals, as a function of the justification given (including no justification) and participant liberalism-conservatism. Ribbons represent 95% confidence intervals. Points are jittered to reduce overlap. Note that the Israel and Conspiracy slopes are largely overlapping.

Does every justification license antisemitism for liberals?

Unlike in Experiments 1 and 2, there were no significant Justification $\times$ Participant Liberalism-Conservatism interactions ($ps > .460$). This means that the effect of each justification (relative to non-justified antisemitism) was not significantly different for liberals vs. conservatives.

The main effects of justification from this model (i.e., the effect at the mean level of liberalism-conservatism), suggested that when individuals justified their antisemitism by disapproval of Israel or by Jewish power and influence, they were liked more (compared to the non-justified expressions of antisemitism): $b = 0.52$, $t(422) = 2.35$, $p = .019$ and $b = 0.53$, $t(422) = 2.42$, $p = .016$, respectively. In contrast, when providing a race-based justification, anti-Jewish targets were liked *less*, $b = -0.54$, $t(422) = -2.52$, $p = .012$. Finally, a religious

justification did not significantly influence liking compared to non-justified antisemitism ($p = .147$).¹

In sum, Experiment 3 partially replicated our findings from Experiments 1 and 2. In those experiments, the effects of the human rights justifications (and of the American values justifications—though with more nuance) were different for liberals and conservatives. When an individual expressing antisemitism justified their prejudice by disapproval of Israel and the war in Gaza or violations of the human rights of Palestinians, liberals, but not conservative, participants liked him or her more.

In Experiment 3, the effect of justifications was not significantly moderated by participant political orientation. Instead, there were some main effects: Americans were more tolerant of expressions of antisemitism when they were justified by disapproval of Israel and the war in Gaza (compared to when no justification was given). Additionally, expressions of antisemitism justified by Jewish power and influence showed a similar licensing effect. However, as these effects are not statistically significant when probed at liberal and conservative levels, we are cautious in interpreting them as definitive.

Candidate Mechanisms

Here, we explore two candidate mechanisms for our pattern of findings: ingroup favoritism and alliance politics. We do so using moderated parallel mediation analyses (Preacher & Hayes, 2008) in R with the *mediation* package (Tingley et al., 2014) and the *emmeans* package (Lenth, 2023) for estimating conditional effects on the A path (i.e., the effect of justification on the candidate mediators).

¹ When we probe the model at both low and high levels of liberalism-conservatism (17 and 81 on the 0-100 liberalism-conservatism scale), none of these main effects remain statistically significant (all $ps > .132$). We note that the effect of the Israel-based justification was in the hypothesized direction, and the same as in Experiments 1 and 2, such that liberals liked the anti-Jewish target using this justification more than when no justification was provided, $b = 0.58$, $t(422) = 1.96$, $p = .050$. This estimate is within the confidence interval of the coefficient in Experiment 2 (0.25 to 1.32), and despite not being statistically significant, contributes additional confidence in the replicated effects in Experiments 1 and 2 (Lakens & Etz, 2017). A Johnson-Neyman test found that the effect of the anti-Israel justification was significant between the values of 18.15 and 61.49; the effect of the conspiracy justification was significant between the values of 27.13 and 63.37; the racist justification had a significant (and negative) effect at values below 64.83; and the religious justification was not significant at any value of the 0-100 liberalism-conservatism scale.

In Experiments 1 and 2, we analyzed each prejudiced target separately using ordinary least squares regression (rather than LMMs). We analyzed only targets where there was a significant effect of justification (i.e., only the anti-Jewish target for Experiment 1).

We tested whether perceived ideological distance and inferred bigotry simultaneously mediated the effects of justifications on liking. “Perceived ideological distance” was computed by rescaling participants self-reported liberalism-conservatism and perceived liberalism-conservatism of each target to a 0-1 scale, then taking the absolute difference between the two. This measure thus ranged from zero (no difference between self-reported liberalism-conservatism and perceived liberalism-conservatism) to 1 (the two were at opposite extremes of the scale).

We also tested whether these indirect effects differed between liberal and conservative participants (defined, as before, as 1 SD below or above the mean liberalism-conservatism). The two mediators were included simultaneously in the model to estimate their unique indirect effects, controlling for each other. We used nonparametric bootstrapping with 1,000 resamples to estimate 95% confidence intervals for indirect effects. There were no missing data.

Statistical Mediation Results

Mediation results for Experiments 1-3 are shown in Figures 4-6. Results from all three experiments suggest that both potential mechanisms contribute to the effects of justifications on liking of prejudiced individuals. We unpack the two mechanisms separately below.

Does Perceived Ingroup Membership Statistically Mediate the Effects of Justifications on Liking? Yes—especially for liberals. As shown in Figures 4-6, in all three studies, expressing antisemitism accompanied by left-coded justifications made liberals view the prejudiced targets as less ideologically distant (i.e., as more ideologically similar) than the targets who didn’t employ justifications, and the indirect effect (Justification → Ideological Distance → Liking) was significant in all three studies. This was not restricted to targets expressing antisemitism. In Experiment 2, liberal participants viewed all prejudiced

targets using human rights justifications as less ideologically distant, and all five prejudiced targets using American values justifications as more ideologically distant. Significant indirect effects suggest that this statistically mediated effects on liking of these targets by liberal participants.

In Experiment 3, individuals expressing antisemitism accompanied by the conspiracy justification were also rated as more ideologically similar by liberals. This was unexpected, as narratives about Jewish power and influence have traditionally come from the political right (and indeed correlate with constructs such as right-wing authoritarianism; Kofta et al., 2020). Finally, liberals viewed the prejudiced target using the religion justification as *more* ideologically distant, and this partly explained why they liked this target less.

The pattern of findings for conservative participants was less clear. Conservatives rated targets employing the left-coded (Israel; human rights) justifications as somewhat more ideologically distant in Experiments 1 and 2, although this effect was not statistically significant in Experiment 3. They rated targets employing other human rights justifications (Experiment 2) as somewhat less ideologically distant, but the effects were smaller than for liberals, and they were not statistically significant in every case.

Does Inferred Bigotry Statistically Mediate Effects of Justifications on Liking? Yes. As shown in Figures 4-6, inferred bigotry statistically mediated the effects of justifications on liking in most cases. In Experiment 1, inferred bigotry had positive indirect effects for both liberals and conservatives. That is, both liberals and conservatives viewed the prejudiced target who expressed disapproval of Israel as less likely to be bigoted, and consequently liked this target more (though the indirect effect was not statistically significant for conservatives).

In Experiment 2, liberal perceivers viewed all targets using the left-coded (human rights) justifications as less likely to be bigoted, and liked these targets more. The indirect effects were again in the same direction for conservative participants but were smaller and not statistically significant in every case. Last, liberals viewed targets employing the right-coded (American values) justifications as more likely to be bigots, and liked these targets

less. This pathway was not statistically significant among conservatives for any of the five prejudiced targets using the American values justifications.

In Experiment 3, inferred bigotry significantly mediated the effect of the justification appealing a disapproval of Israel—liberals viewed these targets as less bigoted, and in turn liked them more. Further, both conservative and liberal perceivers viewed the target using a race-based justification as *more* likely to be bigoted, and in turn liked that target less. Other indirect effects in Experiment 3 were not statistically significant.

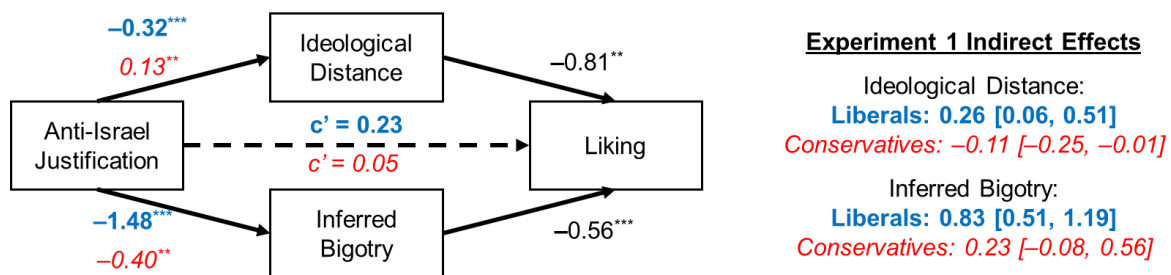


Figure 4. Moderated mediation results from Experiment 1, testing whether the effect of the Anti-Israel Justification (vs. No Justification) on liking is statistically mediated by perceived ideological distance and inferred bigotry. On each path, coefficients for liberals are on top (bolded, in blue) and coefficients for conservatives below (italicized, in red). Coefficients for liberal (blue, bolded) and conservative (red, italicized) perceivers are derived from simple slope estimates at $-1 SD$ and $+1 SD$ from the mean political liberalism-conservatism scale, respectively. Note: * $p < .05$; ** $p < .01$; *** $p < .001$.

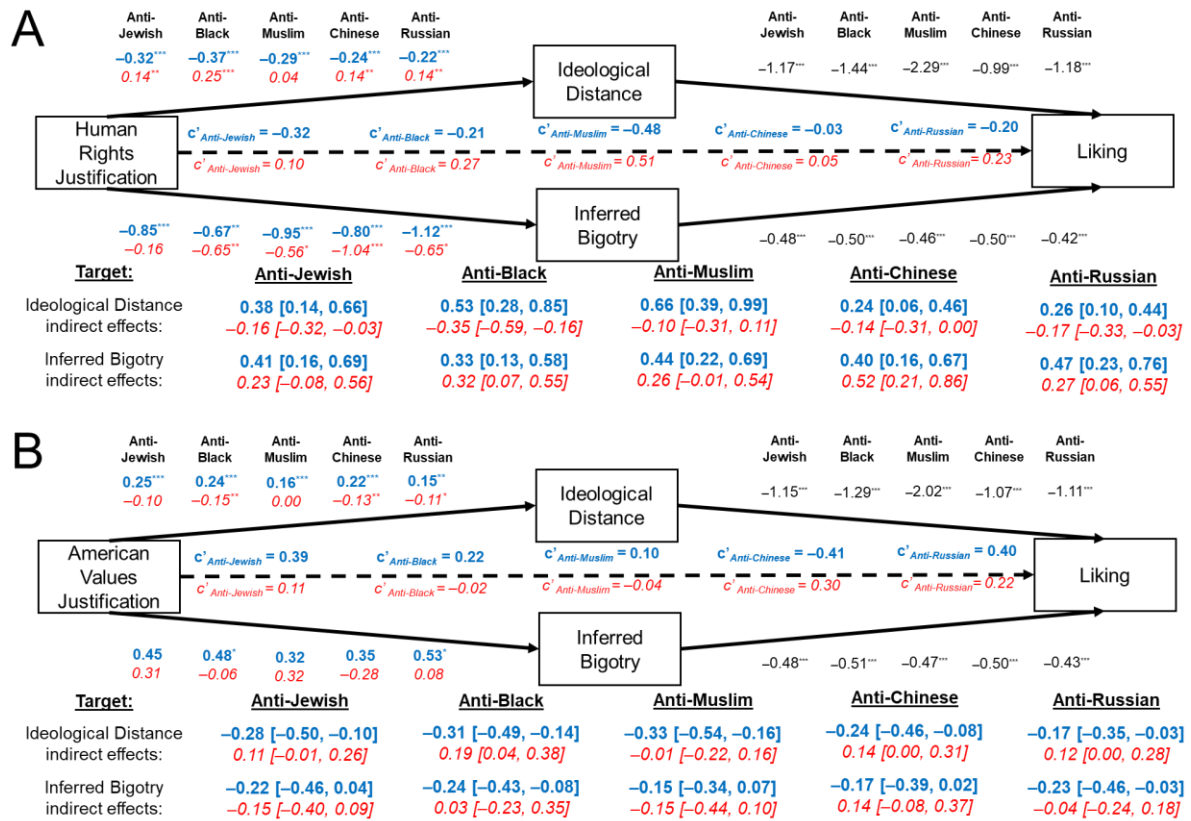


Figure 5. Moderated mediation results from Experiment 2, showing the estimated results of human rights justifications (Panel A) and American values justifications (Panel B) on liking, via ratings of ideological distance and inferred bigotry. Coefficients for liberal (blue, bolded) and conservative (red, italicized) perceivers are derived from simple slope estimates at -1 SD and $+1$ SD from the mean political liberalism-conservatism scale, respectively. Note: $*p < .05$; $**p < .01$; $***p < .001$.

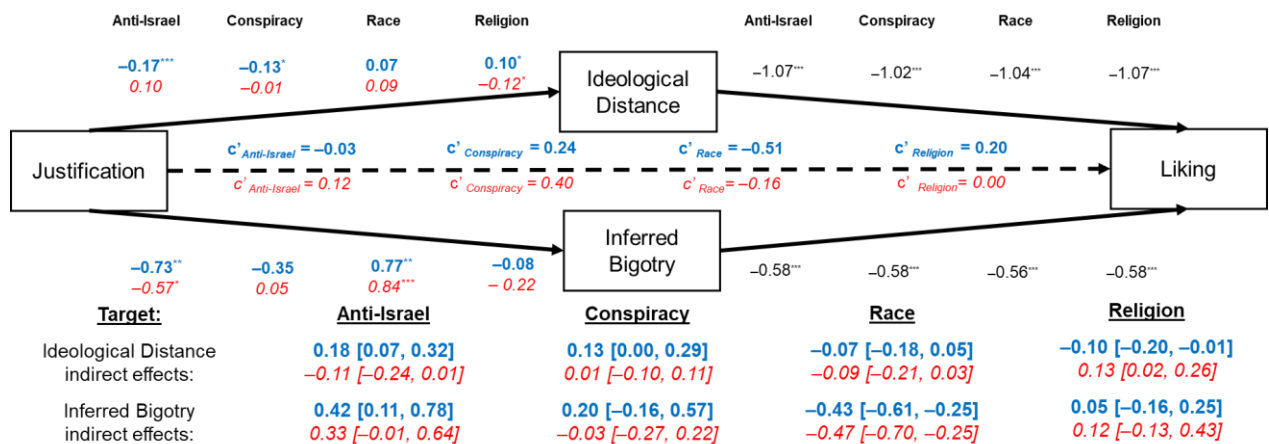


Figure 6. Moderated mediation results from Experiment 3, showing the estimated results of all four justifications (vs. the target with no justification), via ratings of ideological distance and inferred bigotry. Coefficients for liberal (blue, bolded) and conservative (red, italicized) perceivers are derived from simple slope estimates at $-1 SD$ and $+1 SD$ from the mean political liberalism-conservatism scale, respectively. Note: * $p < .05$; ** $p < .01$; *** $p < .001$.

Discussion

Three preregistered experiments examined perceptions of individuals expressing antisemitism (and other prejudices). Liberals viewed antisemitism expressed without justification negatively (and more negatively than conservatives did)—and about as negatively as they viewed expressions of anti-Black and anti-Muslim prejudices. However, when expressions of antisemitism were justified by disapproval of Israel and the war in Gaza or concerns for the human rights of Palestinians—both left-coded justifications—how liberals (but not conservatives) felt about expressions of antisemitism changed: they liked the individuals expressing antisemitism more. Not restricted to expressions of antisemitism, justifications via concerns for human rights made liberals feel more positively toward individuals expressing other prejudices, too.

Although it may seem obvious that liberals will view individuals expressing antisemitism accompanied by appeals to liberal causes more favorably, we note that it does not necessarily follow that people will view these justifications as making anti-Jewish sentiments more palatable. Indeed, the “Black Sheep” effect might predict that liberals would evaluate these individuals *more negatively*, as people often more harshly judge deviant positions expressed by members of their own group (Marques et al., 1988; Pinto et al., 2010). In our experiments, rather than being especially repulsed by expressions of prejudices accompanied by such ingroup markers, the more that liberals perceived individuals expressing prejudices as being in their ideological ingroup, the more they liked them.

Of note, results were less straightforward for conservative participants. Why did right-coded justifications not consistently increase liking among conservatives? Mediation

analyses suggest one possible explanation: often, when conservatives inferred that a prejudiced individual is more ideologically similar to them (which is associated with *higher* liking), they simultaneously viewed this individual as more likely to be bigoted (which is associated with *lower* liking). It seems that these two inferences often cancelled each other out.

Our results suggest that interventions to combat antisemitism should consider not only how people feel about Jews but also the ways in which they justify their feelings. For example, although we do not suggest that liberals are *generally* more lenient toward antisemitism—our findings show that this is not the case—it is clear that certain framings do license some degree of antisemitism, including because these framings are now an ingroup marker among liberals. In fact, a controversial proclamation, such as an expression of prejudice in the service of a social cause, might be an especially strong signal that someone cares greatly about a cause. To the extent that this cause is associated with a specific ingroup, expressing such a view may signal that the speaker is a particularly committed ingroup member (e.g., “I know I shouldn’t say this, but I care so much about the rights of Palestinians that I can’t help but feel some disdain toward Jews”). Thus, curbing antisemitism in liberal circles might require targeted interventions, aimed at recognizing antisemitism when it occurs and reminding liberals that it is wrong regardless of the guises it takes.

Taken as a whole, our findings may explain the highly divergent perspectives on the national conversation about Israel and Palestine. Perhaps the most visible example of this is the pro-Palestine and anti-Israel protests that have swept through the US since the beginning of the Israel-Hamas War. Some have suggested that these protests, during which Israel was strongly condemned, were simply expressions of concerns for the rights of an oppressed group. Others perceived these protests as blatantly tolerating or even spreading antisemitic hatred and tropes, including ones that cast Jews as colonizers and oppressors.

Our results speak to these diverging perceptions: both sides may be correct in some sense. Those who believe that liberal movements condemn antisemitism may be correct in that liberals seem to be highly unaccepting of expressions of antisemitism absent a

justification. However, those who believe that these movements tolerate antisemitism may also be correct in that left-coded justifications *did* make liberals like individuals expressing antisemitism (and other prejudices) more.

Why Justify Prejudice?

Our experiments focus on perceptions of antisemitism, but raise an interesting, more general question about people who feel and express prejudices: Why might they feel the need to justify their prejudices?

Toward managing our reputation in the eyes of others, people aim to present themselves in ways that portray them in a positive light. To be viewed as threatening, and to potentially be socially excluded as a result, carries great costs. Indeed, people who belong to disliked groups appear to have an acute sense both of the threats they are perceived to pose and of the ways in which they can act to effectively counteract those perceptions (Neel et al., 2013). To the extent that prejudiced individuals understand the specific threats they are perceived by others to pose (e.g., as being bigoted), they might strategically employ justifications of the sort “I’m not a bigot, but in the specific case of *this* group...”

This proposal is compatible with the justification-suppression model of prejudice (Crandall & Eshleman, 2003), whereby individuals typically refrain from expressing prejudices except for when a socially acceptable justification allows them to do so. Given that some justifications are viewed more positively for some audiences, one implication is that people strategically adjust their expressions of antisemitism and the justifications they use depending on their audience. Thus, someone wishing to express antisemitism to a liberal audience might find justifications related to the human rights of Palestinians to be more socially palatable, whereas someone wishing to express antisemitism in a different context might find other justifications to be more socially advantageous.

Is Any Justification Acceptable?

A “mindlessness” perspective (Langer, 1989) might suggest that a wide range of justifications would make prejudices more palatable. This clearly was not the case here. Rather, some justifications improved evaluations of individuals expressing antisemitism,

whereas others did not. Importantly, the effectiveness of different justifications depended on the political orientation of perceivers.

Liberals liked anti-Jewish individuals more when they appealed to disapproval of Israel or concerns for the human rights of Palestinians. Whether liberals liked the anti-Jewish targets more when they appealed to conspiracies about Jewish power and influence is less clear, although they did see them as ideologically closer (compared to those expressing antisemitism absent a justification). These same participants, however, strongly disliked targets appealing to Jewish genetic inferiority. The differential effects of the Israel versus race justifications accord with liberal values regarding intergroup power relations and race biology, respectively. In a similar vein of ingroup-relevance, one might expect antisemitism justified by blame for the crucifixion of Jesus Christ to resonate perhaps only for highly religious Christians—which indeed was the case (see Supplemental Materials).

This raises an interesting question: Do effective justifications need to be based on kernels of truth rather than be manufactured out of whole cloth? If so, this might explain why Jews have been accused of being capitalists by communists in the USSR and communists by capitalists in the US. The existence of both Jewish capitalists and Jewish communists lends some minimal credibility to both accusations, which are strategically deployed for different audiences.

Limitations

One important limitation of the present experiments is that, to maintain some degree of ecological validity, the justifications we employed are not fully comparable across prejudiced targets; this concern remained even in Experiment 2, which aimed to more fully control for justifications across prejudiced targets. Although this limits our ability to make direct comparisons about the exact effects of the class of “human rights” justifications applied across groups, we note that such comparisons are not central to our main point—that some justifications decrease the social penalty for antisemitism and other prejudices, especially for liberal perceivers.

Although we have focused here on two mechanisms through which justifications might increase liking—by signalling membership to a perceiver’s coalition or by reducing inferences of dispositional bigotry—future research should explore other mechanisms by which justifications could be effective. For example, effective justifications might communicate other positive information about speakers, such as intelligence or status. Broadly, our findings suggest that justifications function to reveal additional information about speakers and potentially to manage their reputation in the eyes of perceivers. Additionally, as our analyses only showed statistical mediation, future research might use experimental designs to test the causal role of these mediators.

Conclusion

Do liberals tolerate antisemitism? It depends. Whereas liberals strongly dislike blatant expressions of antisemitism (e.g., “I don’t like Jews”), they tolerate it somewhat more when it is accompanied by some—although certainly not all—justifications (e.g., “I don’t like Jews, because I strongly disapprove of Israel and its war in Gaza”). The licensing effect of such justification has at least partially to do with two related inferences liberals make: (1) the individual expressing antisemitism shares political views, and (2) this individual is not generally bigoted (and therefore not a threat to political alliances). A fuller understanding of how the American left responds to antisemitism requires considering how it is expressed, the social information contained in the expression, and the implications of this information for perceivers.

References

- ADL. (2024, September 1). *U.S. antisemitic incidents skyrocketed 360% in aftermath of attack in Israel, according to latest ADL data*. <https://www.adl.org/resources/press-release/us-antisemitic-incidents-skyrocketed-360-aftermath-attack-israel-according>
- Aiken, L. S., & West, S. G. (1991). *Multiple regression: Testing and interpreting interactions*. Sage.
- Bates, D., Mächler, M., Bolker, B., & Walker, S. (2015). Fitting linear mixed-effects models using lme4. *Journal of Statistical Software*, 67(1), 1–48.
<https://doi.org/10.18637/jss.v067.i01>
- Bergh, R., & Brandt, M. J. (2022). Mapping principal dimensions of prejudice in the United States. *Journal of Personality and Social Psychology*, 123(1), 154–173.
<https://doi.org/10.1037/pspi0000360>
- Billig, M. (1988). The notion of “prejudice”: Some rhetorical and ideological aspects. *Text - Interdisciplinary Journal for the Study of Discourse*, 8(1–2), 91–110.
<https://doi.org/10.1515/text.1.1988.8.1-2.91>
- Brenan, M. (2025). *Less than half in U.S. now sympathetic toward Israelis*. Gallup Inc.
<https://news.gallup.com/poll/657404/less-half-sympathetic-toward-israelis.aspx>
- Chenoweth, E., Hammam, S., Pressman, J., & Ulfelder, J. (2024). Protests in the United States on Palestine and Israel, 2023–2024. *Social Movement Studies*, 1–14.
<https://doi.org/10.1080/14742837.2024.2415674>
- Crandall, C. S., & Eshleman, A. (2003). A justification-suppression model of the expression and experience of prejudice. *Psychological Bulletin*, 129(3), 414–446.
<https://doi.org/10.1037/0033-2909.129.3.414>
- Crandall, C. S., Eshleman, A., & O’Brien, L. (2002). Social norms and the expression and suppression of prejudice: The struggle for internalization. *Journal of Personality and Social Psychology*, 82(3), 359–378. <https://doi.org/10.1037/0022-3514.82.3.359>

- Effron, D. A., & Knowles, E. D. (2015). Entitativity and intergroup bias: How belonging to a cohesive group allows people to express their prejudices. *Journal of Personality and Social Psychology*, 108(2), 234–253. <https://doi.org/10.1037/pspa0000020>
- Everett, J. A. C., Caviola, L., Savulescu, J., & Faber, N. S. (2019). Speciesism, generalized prejudice, and perceptions of prejudiced others. *Group Processes & Intergroup Relations*, 22(6), 785–803. <https://doi.org/10.1177/1368430218816962>
- FBI. (2024). *FBI Crime Data Explorer* [Dataset].
<https://cde.ucr.cjis.gov/LATEST/webapp/#/pages/explorer/crime/hate-crime>
- Feldman, N. (2024, February 27). *The New antisemitism*. TIME.
<https://time.com/6763293/antisemitism/>
- Fiske, S. T., Cuddy, A. J. C., Glick, P., & Xu, J. (2002). A model of (often mixed) stereotype content: Competence and warmth respectively follow from perceived status and competition. *Journal of Personality and Social Psychology*, 82(6), 878–902.
<https://doi.org/10.1037/0022-3514.82.6.878>
- Gross, N. (2013). *Why are professors liberal and why do conservatives care?* Harvard University Press.
- Hegarty, P., & Golden, A. M. (2008). Attributional Beliefs About the Controllability of Stigmatized Traits: Antecedents or Justifications of Prejudice?¹. *Journal of Applied Social Psychology*, 38(4), 1023–1044. <https://doi.org/10.1111/j.1559-1816.2008.00337.x>
- Kofta, M., Soral, W., & Bilewicz, M. (2020). What breeds conspiracy antisemitism? The role of political uncontrollability and uncertainty in the belief in Jewish conspiracy. *Journal of Personality and Social Psychology*, 118(5), 900–918.
<https://doi.org/10.1037/pspa0000183>
- Krasner, J. (2024, November 14). Campuses are ground zero in debates about antisemitism – but that’s been true for 100 years. *The Conversation*.
<https://theconversation.com/campuses-are-ground-zero-in-debates-about-antisemitism-but-thats-been-true-for-100-years-239557>

- Kuznetsova, A., Brockhoff, P. B., & Christensen, R. H. B. (2017). lmerTest Package: Tests in Linear Mixed Effects Models. *Journal of Statistical Software*, 82(13).
<https://doi.org/10.18637/jss.v082.i13>
- Lakens, D., & Etz, A. J. (2017). Too true to be bad: When sets of studies with significant and nonsignificant findings are probably true. *Social Psychological and Personality Science*, 8(8), 875–881. <https://doi.org/10.1177/1948550617693058>
- Langer, E. J. (1989). Minding Matters: The Consequences of Mindlessness–Mindfulness. In L. Berkowitz (Ed.), *Advances in Experimental Social Psychology* (Vol. 22, pp. 137–173). Academic Press. [https://doi.org/10.1016/S0065-2601\(08\)60307-X](https://doi.org/10.1016/S0065-2601(08)60307-X)
- Lenth, R. V. (2023). *emmeans: Estimated marginal means, aka seast-square means* [Computer software]. <https://cran.r-project.org/package=emmeans>
- Lewis, B. (2005, December 1). *The new anti-semitism*. The American Scholar.
<https://theamericanscholar.org/the-new-anti-semitism/>
- Long, J. A. (2019). *interactions: Comprehensive, user-friendly toolkit for probing interactions* (Version 1.1.5) [Computer software].
- Mae, L., & Carlston, D. E. (2005). Hoist on your own petard: When prejudiced remarks are recognized and backfire on speakers. *Journal of Experimental Social Psychology*, 41(3), 240–255. <https://doi.org/10.1016/j.jesp.2004.06.011>
- Marques, J. M., Yzerbyt, V. Y., & Leyens, J. (1988). The “Black Sheep Effect”: Extremity of judgments towards ingroup members as a function of group identification. *European Journal of Social Psychology*, 18(1), 1–16. <https://doi.org/10.1002/ejsp.2420180102>
- Neel, R., Neufeld, S. L., & Neuberg, S. L. (2013). Would an obese person whistle Vivaldi? Targets of prejudice self-present to minimize appearance of specific threats. *Psychological Science*, 24(5), 678–687. <https://doi.org/10.1177/0956797612458807>
- New York Times Editorial Board. (2025, June 14). Antisemitism is an urgent problem. Too many people are making excuses. *The New York Times*.
<https://www.nytimes.com/2025/06/14/opinion/antisemitism-jewish-hate.html>

- Norton, M. I., Vandello, J. A., & Darley, J. M. (2004). Casuistry and Social Category Bias. *Journal of Personality and Social Psychology*, 87(6), 817–831.
<https://doi.org/10.1037/0022-3514.87.6.817>
- Pettit, E. (2025, May 21). At UC Berkeley, the faculty asks itself, do our critics have a point? *The Chronicle of Higher Education*. <https://www.chronicle.com/article/at-uc-berkeley-the-faculty-asks-itself-do-our-critics-have-a-point?sra=true>
- Pietraszewski, D., Curry, O. S., Petersen, M. B., Cosmides, L., & Tooby, J. (2015). Constituents of political cognition: Race, party politics, and the alliance detection system. *Cognition*, 140, 24–39. <https://doi.org/10.1016/j.cognition.2015.03.007>
- Pinsof, D., Sears, D. O., & Haselton, M. G. (2023). Strange bedfellows: The Alliance Theory of political belief systems. *Psychological Inquiry*, 34(3), 139–160.
<https://doi.org/10.1080/1047840X.2023.2274433>
- Pinto, I. R., Marques, J. M., Levine, J. M., & Abrams, D. (2010). Membership status and subjective group dynamics: Who triggers the black sheep effect? *Journal of Personality and Social Psychology*, 99(1), 107–119.
<https://doi.org/10.1037/a0018187>
- Preacher, K. J., & Hayes, A. F. (2008). Asymptotic and resampling strategies for assessing and comparing indirect effects in multiple mediator models. *Behavior Research Methods*, 40(3), 879–891. <https://doi.org/10.3758/BRM.40.3.879>
- R Core Team. (2023). *R: A language and environment for statistical computing*. (Version 4.3.2) [Computer software]. R Foundation for Statistical Computing. [//www.R-project.org/](https://www.R-project.org/)
- Rynhold, J. (2020). Democrats' attitudes toward the Israeli-Palestinian conflict. *Middle East Policy*, 27(4), 48–61. <https://doi.org/10.1111/mepo.12526>
- Silver, L. (2024). *Younger Americans stand out in their views of the Israel-Hamas war*. Pew Research Center. <https://www.pewresearch.org/short-reads/2024/04/02/younger-americans-stand-out-in-their-views-of-the-israel-hamas-war/>

- Sommers, S. R., & Norton, M. I. (2006). Lay theories about White racists: What constitutes racism (and what doesn't). *Group Processes & Intergroup Relations*, 9(1), 117–138.
<https://doi.org/10.1177/1368430206059881>
- The Associated Press. (2024, August 14). *UCLA can't allow protesters to block Jewish students from campus, judge rules*. <https://apnews.com/article/ucla-protests-jewish-students-judge-rules-573d3385393b91dae093a8a8f0861431>
- Tingley, D., Yamamoto, T., Hirose, K., Keele, L., & Imai, K. (2014). mediation: R package for causal mediation analysis. *Journal of Statistical Software*, 59(5).
<https://doi.org/10.18637/jss.v059.i05>
- White II, M. H., & Crandall, C. S. (2017). Freedom of racist speech: Ego and expressive threats. *Journal of Personality and Social Psychology*, 113(3), 413–429.
<https://doi.org/10.1037/pspi0000095>
- Wright, G., Volodarsky, S., Hecht, S., & Saxe, L. (2023). *In the shadow of war: Hotspots of antisemitism on US college campuses*. Brandeis University.