

Information Diets Are More Diverse in Attention Than in Engagement

Christopher Barrie,^a Aybuke Atalay,^b Alia ElKattan^a

a) New York University; b) University of Edinburgh

Abstract: What political content do we pay attention to online? Diverse political information is essential for democratic competence, yet online media raises concerns about fragmented information diets. Research on selective exposure highlights how social media can foster ideological echo chambers, while other studies emphasize incidental exposure to diverse viewpoints. A critical measurement challenge is that public platform traces often observe engagement (e.g., likes or shares) more readily than lower-visibility forms of attention—i.e., what users notice or choose to read without necessarily interacting publicly. In this study, we address this gap with a social media clone platform that separately records attention and engagement. We ran the study in both the United Kingdom (once) and the United States (thrice). Across both contexts, we found that the ideology-engagement association is significantly larger than the ideology-attention association. This underscores the importance of measuring attention, rather than solely engagement, to accurately assess the diversity of online information diets.

Keywords: selective exposure; attention; engagement; social media; political communication; echo chambers

Reproducibility Package: Data, code, documentation, and replication materials necessary to reproduce the empirical and simulation results reported in this article are available at <https://github.com/cjbarrie/smclone>.

Citation: Barrie, Christopher, Aybuke Atalay, and Alia ElKattan. 2026. "Information Diets Are More Diverse in Attention Than in Engagement" *Sociological Science* 13: 864-883.

Received: April 9, 2026

Accepted: June 8, 2026

Published: July 21, 2026

Editor(s): Ari Adut, Kieran Healy

DOI: 10.15195/v13.a33

Copyright: © 2026 The Author(s). This open-access article has been published under a Creative Commons Attribution License, which allows unrestricted use, distribution and reproduction, in any form, as long as the original author and source have been credited. 

THE diversity of information citizens consume shapes democratic discourse and “democratic competence” (Fishkin 1991; Guess 2021; Gutmann and Thompson 2004). The rise of online media—combining personalized algorithms with self-selected networks—has raised concerns about increasingly balkanized information diets. This debate has evolved in phases. Early accounts warned of “echo chambers,” in which algorithmic filtering and homophily produce ideologically congruent feeds (Cinelli et al. 2021; Pariser 2011). Subsequent work complicated this view. Although users do prefer congenial content, many studies show that they also encounter counter-attitudinal information incidentally, through material that appears in their existing networks (e.g., a friend’s share or reply) rather than through content they actively seek out (Bail 2021; Flaxman, Goel, and Rao 2016; Fletcher and Nielsen 2018).

A newer dynamic is out-of-network algorithmic exposure: platforms increasingly inject content from accounts a user does not follow, based on engagement/attention signals and policy choices. Recent governance changes and moderation shifts—especially on X—have altered recommendation priorities and content visibility, affecting who sees what regardless of follow networks (Barrie 2023; Hickey et al., 2023; Ye, Luceri, and Ferrara 2024). This mechanism differs from network-mediated

incidental exposure: it is produced by ranking systems rather than a user's social graph, and it can systematically amplify specific political viewpoints at scale.

Do people pay attention to such content? The answer depends on how attention is measured. Experimental work on selective exposure and selective attention shows that people often choose congenial political information when given the opportunity to select what to read (Feldman et al. 2013; Garrett and Stroud 2014; Graf and Aday 2008; Stroud 2017). At the same time, many studies of large online platforms rely on public engagement (likes/shares), platform-provided exposure measures, or self-reports, which can leave unclear what users deliberately noticed or read without interacting (Lazer 2020). By *engagement*, we mean a visible interaction with a post (liking or sharing), which typically communicates endorsement or affiliation (though not exclusively). By *attention*, we mean an intentional act to access more of the post's informational content, regardless of whether any outward engagement follows. The distinction matters even more as platforms move toward attention-optimized recommendation, including short-video and out-of-network feed systems that rely heavily on implicit behavioral signals such as watch time, dwell time, and repeated viewing (Narayanan 2023).

In this contribution, we are able to separate attention from engagement in a controlled social-media clone environment. Participants scroll a curated, ideologically balanced feed; every post is truncated with a "Read more" expander. We record (1) attention as clicking "Read more" and (2) engagement as likes/shares. The content set, order, and endorsement cues are controlled, allowing clean comparisons across behaviors and studies. Anticipating our results, we find that the ideology-*engagement* association is substantially larger than the ideology-*attention* association across one U.K. and three U.S. studies. Users attend to a wider range of political perspectives than their outward engagements would suggest. We open source the clone environment that we developed for use by future researchers.

Exposure to Politics Online

Ideological diversity is central to democratic competence—the capacity of citizens to make informed, reasoned choices (Downs 1957; Fishkin 1991). Digital media raised concerns that personalization and self-selection would narrow information diets through selective exposure, which refers to the tendency to seek congenial content and avoid dissonant content (Cinelli et al. 2021; Pariser 2011). Such patterns, attributed to cognitive dissonance reduction, confirmation bias, and perceived source credibility, can generate echo chambers or filter bubbles and may heighten affective polarization (Barberá, 2015; Hobolt, Lawall, and Tilley, 2024; Metzger, Hartsell, and Flanagan 2020).

However, complete isolation from opposing views is uncommon (Bail et al. 2018; Guess 2021). A large literature documents *incidental exposure*—encounters with counter-attitudinal content that occur not by active search but via one's network (e.g., a friend's share, reply etc.) (Flaxman et al. 2016; Fletcher and Nielsen 2018). Even with personalization, algorithmic feeds and social sharing regularly introduce cross-cutting material.

Of course, attention is not absent from this literature. Experimental work on selective exposure and selective attention shows that people often choose or attend

to attitude-consistent political information, especially in news-choice settings where political cues are relatively clear (Feldman et al. 2013; Garrett and Stroud 2014; Graf and Aday 2008; Stroud 2017). Less is known about how strongly the same alignment appears for short, social-media-style posts, which users may encounter in feeds where political content is cross-cutting, incidental, weakly cued, or intermixed with other material.

This distinction matters as a newer channel of exposure becomes more prominent: *algorithmic out-of-network exposure*. Ranking systems increasingly inject content from accounts a user does not follow, guided by attention/engagement signals and policy choices. Recent governance and moderation shifts—especially on X—have altered recommendation priorities and the visibility of political viewpoints, reshaping who sees what regardless of follow networks (Barrie, 2023; Hickey et al. 2023; Huszár et al. 2022; Ye et al. 2024). This mechanism differs from network-mediated incidental exposure because it is produced by the recommender itself and can systematically amplify particular perspectives at scale. The broader shift toward recommendation-centric feeds is not limited to X: platforms such as TikTok make out-of-network recommendation a central organizing principle, and recommender systems increasingly learn from implicit attention signals as well as explicit engagement (Narayanan 2023).

Taken together, we can say that exposure to political information might result from three processes: (1) *selective exposure* (following/curating congenial sources); (2) *network-mediated incidental exposure* (cross-cutting items arriving via one's social ties); and (3) *algorithmic out-of-network exposure* (items inserted by the recommender, independent of the social graph).

Accurately assessing the diversity of what users *consume* therefore requires control over three elements: (1) the *content set* users can potentially see; (2) the *exposure mechanism* (ordering/ranking independent of user networks); and (3) a way to record *attention*, not only visible engagement. Our design satisfies these conditions in a social-media clone by curating the content set, randomizing the feed order, and separately logging attention (“Read more”) and engagement (Like/Share), allowing clean comparisons across behaviors.

Measuring Attention

Measuring human attention is a persistent challenge in computational social science (Lazer 2020; Lazer et al. 2021). Communication research emphasizes that attention is not binary but graded, ranging from fleeting exposure to sustained, effortful processing and comprehension (De Vreese and Neijens 2016). Treating mere presence in front of content as “attention” collapses these levels and risks conflating potential exposure with deliberate uptake.

We do not claim to solve this thorny measurement problem. Our contribution, instead, is a pragmatic advance for studying online information diets: we operationalize *attention* as a click on a “Read more” expander that reveals the remainder of an otherwise truncated post (first ~100 characters). In our framework, a “Read more” click is therefore evidence of intentional, content-seeking attention; whereas *engagement* refers to visible interactions (likes/shares).¹

To date, researchers have pursued multiple approaches, each with distinct advantages and limitations, which we lay out below.

Engagement

Digital-trace approaches infer political exposure and preferences from follows and public, endorsement-like interactions (e.g., likes and reshares) on social platforms. Analyses of networks and interaction patterns show that online political talk is homophilous, with users clustering by ideology (Barberá et al. 2015; Cinelli et al. 2021). However, these traces primarily record *engagement*, not *attention*: they capture visible actions while omitting content that users notice or read without interacting (Lazer 2020). This construct mismatch can bias inferences about information diets by undercounting cross-cutting exposure that does not manifest as public engagement.

Studies with privileged access to feed exposures and click logs quantify ideological segregation along an exposure-engagement gradient in the United States (Bakshy, Messing, and Adamic 2015; González-Bailón et al., 2023; Nyhan et al. 2023). Notably, González-Bailón et al. (2023) find that segregation is strongest for active engagement and attenuates toward actual or potential exposure. One of our contributions is to isolate this gradient *without* platform access: we fix the content set, randomize its order for each respondent, and directly observe attention via a “Read more” click in a controlled clone environment, enabling a clean comparison with engagement.

Surveys

Another common approach relies on surveys, including so-called “linked surveys.” Standard survey methods ask participants directly about their media habits to estimate information diversity (Dubois and Blank 2018). However, self-reported data suffer from significant reliability issues: respondents often inaccurately recall or intentionally misrepresent their media consumption, thereby measuring reported rather than actual attention (Guess 2015).

To mitigate these limitations, researchers have used linked surveys, combining self-reported data with direct observations from respondents’ social media accounts (Guess 2021; Guess et al. 2019). By incorporating data such as follows, likes, and shares, as well as web browsing data, this method better captures overall information exposure. Nonetheless, even this improved approach misses important passive attention to, for example, information incidentally or algorithmically injected into a user’s feed. In addition, even with the complete follow network, closed-source algorithms and the inability to measure actual attention with current linked survey designs mean we omit a large portion of the information we actually consume online (Bakshy et al. 2015; Lazer et al. 2021). Finally, linked surveys are resource-intensive, restricting their scalability for broader studies.

Eye-Tracking

A third approach uses eye-tracking and related experimental measures to capture moment-to-moment visual attention to on-screen content. This method offers fine

temporal resolution and spatial precision (fixations and saccades) (Bode, Vraga, and Troller-Renfree 2017). Some selective-attention experiments find that people attend more to attitude-consistent political information (Graf and Aday 2008; Feldman et al. 2013; Stroud 2017); other eye-tracking evidence suggests that ideological alignment is a stronger predictor of *engagement* than of *visual attention* per se (Sülflow, Schäfer, and Winter 2019). Taken together, these results underscore a key point: looking, reading, expanding, endorsing, and sharing are distinct processes, and engagement metrics alone can miss substantial passive consumption.

At the same time, eye-tracking poses logistical and inferential challenges. It typically requires specialized hardware and controlled settings, which limits ecological validity and sample size (Duchowski 2007). These constraints call for complementary approaches that can be deployed at scale in realistic settings. Our design trades sub-second gaze data for a behaviorally verified, interpretable proxy of attention: clicking “Read more” to reveal truncated text. By holding the content set fixed, randomizing the order of posts for each respondent, and logging both attention (expansions) and engagement (likes/shares), we create a clean separation between the two types of behavior.

Data and Method

We study behavior in a custom-built social-media clone environment that gives us full control over what respondents could see and how it was shown. The platform (1) fixes the content set and its ideological balance, (2) randomizes timeline order for each respondent, and (3) separately logs *attention* (a “Read more” expansion of truncated text) and *engagement* (Like/Share clicks). We also randomize the displayed endorsement (likes+shares) in calibrated bins to isolate visual popularity cues from content. These randomizations are not experimental manipulations: our estimands are *associations* between respondent ideology and behavior under a standardized presentation regime. We do not identify causal effects of ideology or content; rather, we compare how the ideology-behavior gradient differs for attention versus engagement when exposure opportunities and endorsement signals are held exogenous to user or tweet characteristics.²

We ran four online studies with participants from the United Kingdom and the United States, each using a custom-built clone social media environment tailored to the political context of that country. Although the design was held constant across countries, the tweets presented to users differed across countries in order to better represent domestic political information environments in the United Kingdom and United States.

Candidate Tweet Collection

UK Candidate Tweets. For the U.K. context, in order to seed a sample of political tweets, we relied on a dataset of 2024 U.K. General Election candidate accounts provided in Barrie, Atalay, and ElKattan (2025a). This yielded 2,729 unique X usernames. We then took the individual X user IDs for each candidate and used these to gather the last 3,200 posts for each user with the `rettiwt` JavaScript library

(Rettigt API Contributors 2023). We began collection on June 12, 2024. We were ultimately able to gather data for 1,603 users. We then filtered the data to include only posts from May 22, 2023 to May 22, 2024—that is, the one-year period before the snap General Election announcement. We also filtered out any retweets or replies to other posts before taking a subsample of 10k posts.

U.S. Senator Tweets. For the U.S. context, we generated a sample of politically relevant tweets by relying on data collected by the Center for Social Media and Politics at New York University. Specifically, we used a regularly updated list of current U.S. Senators as of 2020. This list was initially compiled by taking the names of U.S. Senators from the official roster and then locating their verified X handles. Once the X IDs were confirmed (via the account's verified badge and profile information), the collection of tweets from these accounts was automatically updated thrice per week (every Monday, Wednesday, and Friday). The underlying list of senators is itself updated in each federal election cycle (e.g., 2020, 2022, and 2023), ensuring that newly sworn-in or departing senators are added or removed before the next round of harvesting. To seed our annotation pipeline, we again took a 10k post subsample.

Annotation and Sampling. After collecting tweets from political candidates in each country, we annotated their ideological orientation. To classify tweets ideologically, we used the OpenAI language model gpt-4o-mini-2024-07-18 (accessed via the OpenAI API, June–July 2024). The model was prompted to assign tweets to a five-point scale ranging from “Very left” to “Very right” based on the political opinion expressed. The same classification procedure was applied to both the U.K. and U.S. datasets to ensure consistency across contexts. Model classifications were subsequently reviewed by the authors, and labels were retained only when a majority agreement was reached between the model output and human coders. Finally, we generated a balanced sample across annotations by slicing tweets into equal-sized bins from Very left to Very right. Specifically, we randomly sampled from each ideological category down to the size of the smallest group, ensuring all bins contained an equal number of tweets.

Importantly, when annotating the data with gpt-4o-mini-2024-07-18, we deliberately truncated the text to be classified after 100 characters at which point the text then read “...Read more”. We did this in order to effectively recreate exactly how the tweet would appear to users of the clone social media environment. Because ideological classification was performed on the truncated text rather than the full tweet, the ideological label reflects the portion of the post visible to participants prior to clicking “...Read more.” If it contained no political opinion, the model was prompted to classify as “Neither left nor right.” We print the full prompt design in the online supplement. The identical truncation rule was applied in both countries. In the online supplement, we report the precise model version and prompt format, which supports re-use of our procedure. We nonetheless acknowledge that hosted models can change over time (Barrie, Palmer, and Spirling 2025b).

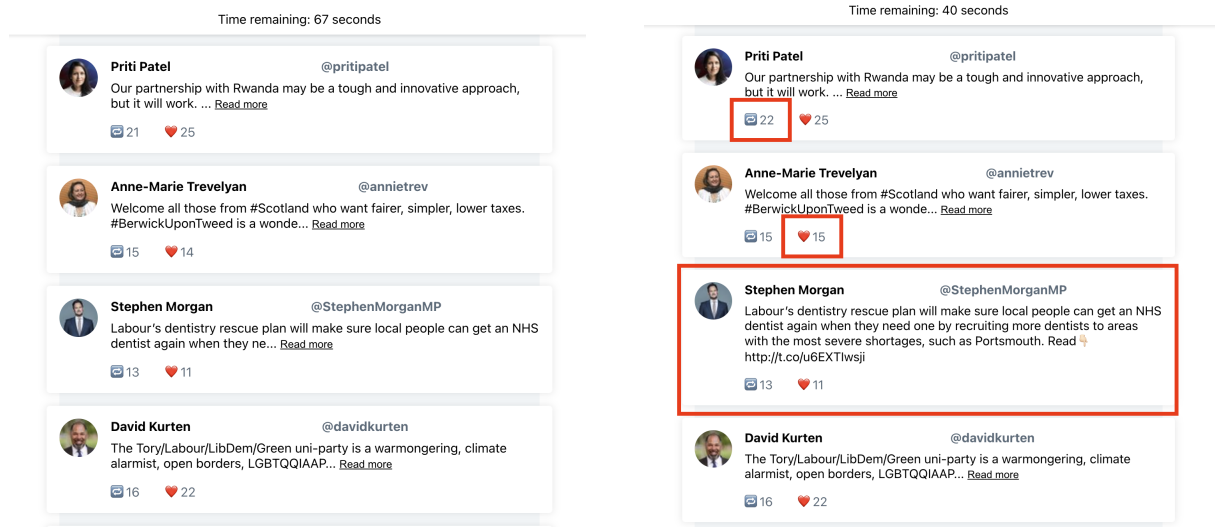
All three authors then inspected the annotations and corrected any for which there was a common disagreement in both sets of tweets. That is, each author

inspected the tweet and classified it from “Very left” to “Very right”. We kept only those for which there was majority agreement between the authors and the large language model (LLM). We kept codebook guidance for authors on how to judge the political cue of a tweet at a minimum. We did this deliberately because we wanted to record the immediate, instinctive, reaction of a human user to a piece of political content—similar to how they would judge the content in the real world. That is, we did not want the authors to annotate tweets according to some codebook criteria but instead react to what kind of political outlook they thought it indicated. We also removed tweets that obviously referenced a specific time period that predated the experiment—for example, posts referring to Rishi Sunak as the current Prime Minister in the United Kingdom or referencing Joe Biden as the current President of the United States. We similarly excluded tweets tied to specific calendar events, such as New Year celebrations, that could signal the timing of the original post.³ Finally, we constructed country-specific balanced feeds: the U.K. study used 25 tweets from each of the four classifications “Very left,” “Somewhat left,” “Somewhat right,” and “Very right” (100 posts total), whereas the U.S. studies used 30 tweets from each classification (120 posts total).

This truncation procedure was also intended to address a visibility concern: the ideological label should be based on the same cue initially available to respondents, not on information revealed only after a “Read more” click. Still, the procedure does not guarantee that every participant perceived the political cue in the way the LLM or authors did. Some posts may have required more political knowledge, more careful reading, or more contextual inference than others. This is an important limitation for external validity because the ideological signal in a real platform feed may be more or less obvious depending on the source cue, image, algorithmic context, and surrounding posts. If ideological cues are more obvious in real platform environments than in our clone, the downstream expectation would be a more pronounced ideology-attention gradient than the one we estimate here, as users may avoid expanding obviously counter-attitudinal content.

To further assuage concerns about measurement error at the level of tweet ideological slant, we go on to simulate a series of error regimes as a means of determining the robustness of our results to different types of potential mislabeling, including a worst-case procedure that deliberately flips the labels most favorable to our conclusions.

Study Procedure and Measures. We embedded each post in a custom social-media clone. Upon arrival, respondents consented and were told that the site is a research platform where they could interact by clicking “Read more,” “Like,” or “Share.”⁴ Each participant then viewed the timeline for 120 seconds and could interact freely. The posts were randomly ordered for each respondent: 100 posts in the U.K. study and 120 posts in each U.S. study. This removes person-specific ordering as a confound and attenuates position and recency effects by randomizing for each user. For every post, the displayed likes and retweets were independently randomized by drawing from bins designed to approximate the observed empirical distribution, providing a social-endorsement signal that is orthogonal to content and respondent traits (Messing and Westwood 2014).



(a) Example clone environment with which respondents interacted.

(b) Example of post interactions (Share, Like, and Read more).

Figure 1: Clone environment.

We integrated posthog-js within the platform to stream granular click events to PostHog (Team 2024). For each respondent-post pair (i, j) we record: (1) “Read more” expansions (our operationalization of *attention*) and (2) “Like”/“Share” clicks (our *engagement* measures).⁵ We provide an example of the clone environment respondents saw below in Figures 1a and 1b.

After the interaction window, respondents completed a short questionnaire. Our key predictor is a standardized ideology (0–10 self-placement) measure, which asked: “In political matters, people talk of ‘the left’ and ‘the right.’ How would you place your views on this scale, generally speaking?”. We also asked participants questions about political interest (five-point Likert), and standard demographics (age, gender, and education; income in the United Kingdom only), plus self-reported social-media use. All additional item wordings and motivations are in the online supplement, alongside details of the attention checks we included for each country.⁶ Within each country, we applied the same content curation rules and randomization scheme (the three U.S. studies share an identical protocol, and the U.K. study uses an analogous, country-specific protocol).

Hypotheses and Estimation Strategy

Building on findings from previous work, we preregistered two core hypotheses (see preregistration: <https://osf.io/ny9du/overview>).

1. **H1:** Social media users are no more or less likely to pay attention to content that aligns with their own political orientation.⁷
2. **H2:** Social media users are more likely to “like” or “retweet” content that aligns with their own political orientation.

Hypothesis H1 aligns with research finding little or no association between a user's ideology and the ideological content of material they visually attend to (Sülflow et al. 2019), but it sits in tension with other selective-attention and selective-exposure experiments that find congeniality effects in attention or consumption (Feldman et al. 2013; Garrett and Stroud 2014; Graf and Aday 2008). Related platform-trace evidence shows that ideological segregation is strongest for active engagement and attenuates as the measure moves toward actual or potential exposure (González-Bailón et al. 2023). We therefore treat H1 cautiously: the central substantive question is not whether selective attention is entirely absent, especially for conventional news media choices where ideological alignment is well documented (Stroud 2010), but whether the ideology gradient for attention remains small relative to the ideology gradient for visible engagement in a feed of short, social-media-style posts that users may encounter incidentally. Hypothesis H2 follows work documenting a clear preference for ideologically congruent information in measures of engagement and following behavior (Barberá 2015; Dubois and Blank 2018).⁸

Taken together, H1 and H2 imply a more precise contrast that is central to our design: the *ideology gradient should be steeper for engagement than for attention*. The estimation strategy that follows is therefore structured around two related goals:

- (1) estimating the ideology-congruence gradient $g^{(k)}$ for each interaction type within each study, and
- (2) directly estimating the within-study contrasts $\Delta^{(k)}$.

We follow Lundberg et al. (2021) in first defining our theoretical and empirical estimands, target population, and estimation procedure.

Target Population. For the U.K. study, our target population is defined as individuals residing in the United Kingdom, ranging from ages 18–65, both male and female, who have some knowledge of the Internet (given we used online panels). Participants for the U.K. experiment were recruited through the Prolific platform.

For the U.S. studies, the target population is similarly defined as individuals, aged 18–65, who reside in the United States and have some knowledge of the Internet. Participants were recruited through two sources: Prolific and RIWI. Although Prolific draws from a pool of preregistered online panelists, RIWI uses its proprietary “Random Domain Intercept Technology,” which invites users to surveys via pop-ups triggered when they click on a broken link. The RIWI sample is filtered by geolocation and age, including only users located in the United States and excluding anyone who self-reports being under 18.⁹ The Prolific sample recruits from their “Standard sample” of U.S. residents. For the first U.K. and first U.S. Prolific experiments, we conducted small pilots to ensure the platform was functioning as expected. We excluded participants in the pilot from the main experiment.

In addition to a conventional convenience panel, we fielded a Prolific Representative study in the United States. Although not a probability sample, this sample stratifies recruitment to approximate national benchmarks on key demographics, which mitigates some concerns about the generalizability of convenience samples.¹⁰

We use this study to assess whether the descriptive associations we document are robust to a sampling frame with improved population coverage.

Estimands. In our design, every respondent is assigned the same fixed set of tweets in a randomly ordered feed. For each respondent i , tweet j , and interaction type $k \in \{\text{Read, Like, Share}\}$, let

$$\ell_{ij}^{(k)} = \text{logit } \Pr\{A_{ij}^{(k)} = 1 \mid \text{assigned tweet } j\},$$

where $A_{ij}^{(k)}$ is an indicator that respondent i clicked “Read more,” “Like,” or “Share” on tweet j at least once. We define the estimand on the log-odds scale because the empirical model below is a fixed-effects logistic regression.

Let P_i denote standardized ideology (mean 0, SD 1) and $D_j \in \{0, 1\}$ an indicator that tweet j is right-aligned (1) versus left-aligned (0). This binary indicator collapses the four tweet-label categories used in sampling: “Very left” and “Somewhat left” are coded 0, while “Somewhat right” and “Very right” are coded 1. For a given interaction type k , our target quantity is the *congruence gradient*: how much ideology changes the difference in log-odds of acting on right- versus left-aligned tweets, conditional on respondent and tweet fixed effects. That is, for each k we want

$$g^{(k)} = \text{change in } [\ell^{(k)}(P_i, D_j = 1) - \ell^{(k)}(P_i, D_j = 0)]$$

when P_i increases by 1 SD. This answers the following question: as respondents become more right-leaning by one standard deviation, how much more strongly are they associated with acting on right-aligned rather than left-aligned tweets, conditional on respondent and tweet fixed effects.

We compute $g^{(k)}$ separately for $k = \text{Read, Like, Share}$. Our main comparison is between attention and each engagement behavior in turn. Taking “Read more” as the attention baseline, we define for each engagement type $k \in \{\text{Like, Share}\}$ a separate contrast

$$\Delta^{(k)} = g^{(k)} - g^{(\text{Read})},$$

which tells us, for that specific engagement behavior k , how much more strongly ideology predicts partisan-congruent actions compared to read-more clicks.

Estimation Strategy. For each study, we build a respondent-tweet exposure panel. We are able to do this because all respondents in a given study were assigned the same tweet set.¹¹ We define the baseline analysis sample as respondents who (1) provided a valid participant ID in the survey and (2) generated at least one outcome event in the PostHog logs (e.g., a click on “Like”). For every such respondent-tweet pair (i, j) in the standard tweet set, we create one row and define three binary outcomes

$$A_{ij}^{(\text{Read})}, A_{ij}^{(\text{Like})}, A_{ij}^{(\text{Share})} \in \{0, 1\},$$

each equal to 1 if respondent i ever clicked that action on tweet j , and 0 otherwise. We then stack these outcomes into long format with an event-type indicator $k \in$

{Read, Like, Share}, and merge in standardized ideology P_i and tweet orientation D_j .

We estimate, separately for each study, a single joint fixed-effects logistic regression on this long panel. Let $A_{ij}^{(k)}$ denote the stacked outcome for action type k on tweet j by respondent i , P_i standardized ideology (mean 0, standard deviation 1), and $D_j \in \{0, 1\}$ an indicator that tweet j is right- (vs. left-) aligned. The implemented model has action-specific intercepts and action-specific ideology-by-orientation slopes, with respondent and tweet fixed effects held common across action types. It can be written as

$$\text{logit } \Pr(A_{ij}^{(k)} = 1 \mid P_i, D_j, k, \alpha_i, \gamma_j) = \beta_0^{(k)} + \beta_3^{(k)}(P_i D_j) + \alpha_i + \gamma_j,$$

implemented in practice with a `fixest` formula that interacts the ideology-by-orientation term with event type and includes tweet and respondent fixed effects, with “Read more” as the event-type baseline. The respondent fixed effects α_i and tweet fixed effects γ_j are not interacted with action type; they are common fixed effects in the stacked model. This is intentional: the model allows the baseline action rate and ideology-by-orientation gradient to differ across Read, Like, and Share, while conditioning on a common respondent and tweet baseline when comparing action types.

For “Read more,” the coefficient

$$g^{(\text{Read})} = \beta_3^{(\text{Read})}$$

is the *congruence gradient* for the attention outcome: the change in the log-odds of taking action on a right- versus a left-aligned tweet associated with a one-standard-deviation shift in respondent ideology, conditional on respondent and tweet fixed effects. The reported Like and Share rows in the regression tables are incremental differences relative to Read more, corresponding to $\Delta^{(\text{Like})}$ and $\Delta^{(\text{Share})}$. The full Like and Share gradients are therefore obtained by adding the reported Like–Read or Share–Read delta to the reported Read-more gradient. Because everyone is assigned the same tweet set and we include all respondent-tweet pairs for the analysis sample, these gradients are interpretable as ideology-by-orientation slopes in the log-odds of clicking a given tweet, conditional on assignment to the feed.

For $k \in \{\text{Like}, \text{Share}\}$, we then estimate the linear contrast of the corresponding interaction coefficients in the joint model. Positive values of $\Delta^{(k)}$ indicate that ideology is more strongly associated with partisan congruence in likes or shares than in read-more clicks.

We compute study-specific standard errors and confidence intervals from the model variance-covariance matrix with two-way clustering by respondent and tweet. Because ideology is standardized within each study, all gradients and differences can be interpreted as changes in log-odds for a one-standard-deviation shift in respondent ideology; exponentiating these coefficients yields the corresponding odds ratios.

To summarize results across studies, we meta-analyze the study-specific gradients and differences using both a common-effect (“fixed-effect”) model and a random-effects model estimated by restricted maximum likelihood (REML). The

common-effect specification assumes a single underlying effect and differences across studies arising only from sampling error (with inverse-variance weighting), whereas the random-effects specification allows for genuine between-study heterogeneity in the underlying gradients.¹² The random-effects specification does not assume a single common effect; instead, it allows each study to have its own underlying gradient and estimates the mean of this distribution alongside a between-study variance parameter. We therefore interpret the pooled estimates as average congruence gradients across these four implementations of the clone environment, rather than as a universal effect size for all platforms, time periods, or tweet universes.

As a robustness check, we also estimate an expanded respondent-tweet joint fixed-effects model. In contrast to the baseline clicker-restricted assigned-post panel, this specification begins from the set of respondents who logged into the interface with a valid participant ID and assigns $A_{ij}^{(k)} = 0$ when no click is recorded at all.¹³

Results

The results, including meta-analytic estimates, are displayed in Figure 2. All models use a standardized ideology measure. Within each study, we rescale respondents' left-right self-placement to have mean 0 and standard deviation (SD) 1, so a coefficient corresponds to the effect of moving one standard deviation to the right in that study's ideology distribution. The joint models are logistic regressions, so coefficients are on the log-odds scale and exponentiating gives the odds ratio: the factor by which the odds of acting on a right- (vs. left-) aligned tweet change for a one-SD shift in ideology.

We first consider "Read more" clicks, which we understand as a proxy for attention. Across the four studies, the ideology gradient for read-more clicks is small. Study-specific estimates center near zero (e.g., 0.08 in the United Kingdom; -0.06, 0.10, and 0.20 in three U.S. studies), and the meta-analytic estimate is about 0.07, corresponding to an odds ratio of $\exp(0.07) \approx 1.08$. Thus, moving one SD to the right in ideology increases the odds of expanding a right- versus a left-aligned tweet by only about 8percent; attention is only weakly sorted along partisan lines. We do not have support for there being *no effect* (H1), though we can say that the substantive size of this association is small.

Likes and shares display much steeper ideology gradients. In the U.K. study, the gradients are about 0.56 for likes and 0.60 for shares (odds ratios ≈ 1.76 and 1.82); a one-SD rightward shift raises the odds of liking or sharing a right-aligned tweet by roughly 75–80percent. In later U.S. studies, the gradients are larger still: for example, in the representative U.S. sample the like and share gradients are 1.14 and 1.28 (odds ratios ≈ 3.13 and 3.60), respectively. The pooled meta-analytic gradients are ≈ 0.63 for both likes and shares, implying odds ratios of $\exp(0.63) \approx 1.9$. On average, a one-SD move to the right almost doubles the odds of engaging with congenial content.

To compare attention and engagement directly, we subtract the read-more gradient from each engagement gradient within each study. These differences are positive in every study and typically large. In the U.K. study, the like-read and

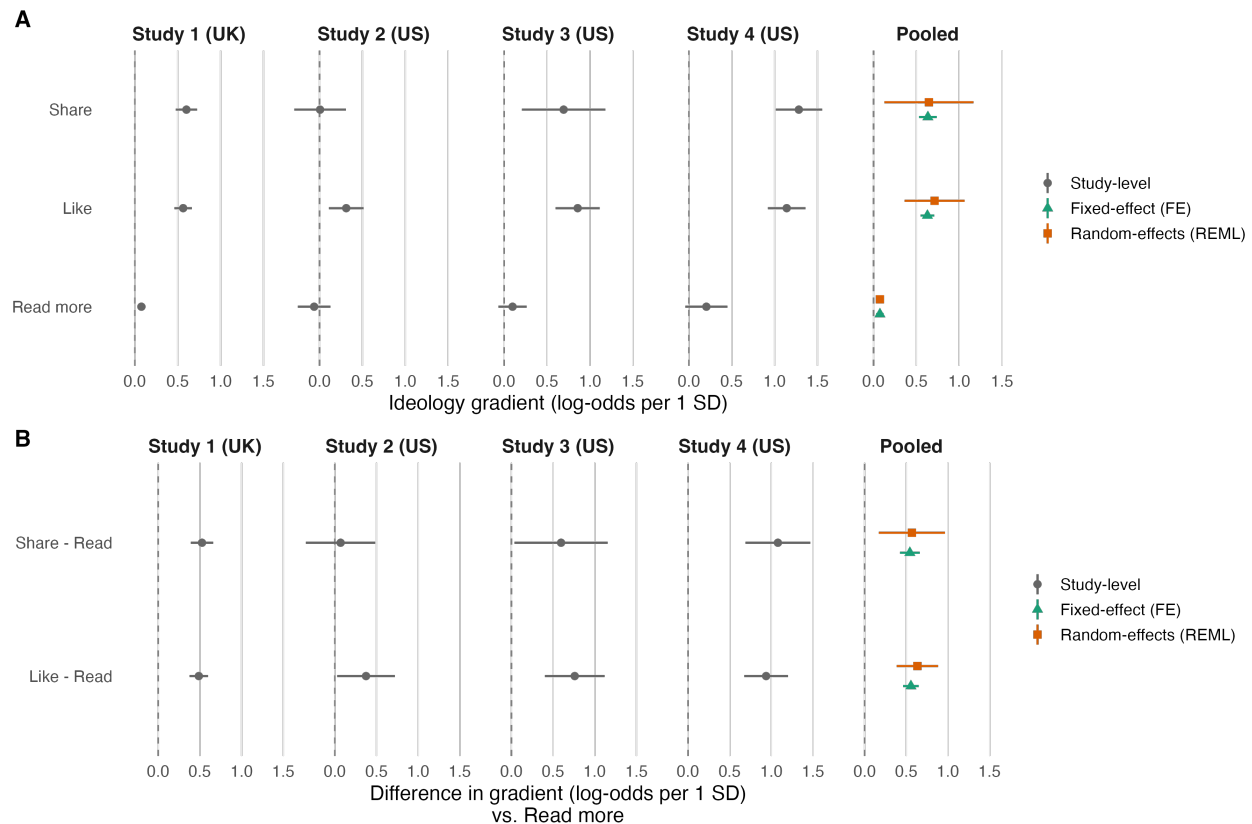


Figure 2: Joint fixed-effects estimates of ideology gradients across four studies and pooled meta-analytic models. The top row shows, for each study and for the pooled estimates, the ideology–congruence gradient for each interaction type (“Read more”, “Like”, “Share”): the change in log-odds of acting on a right- versus left-aligned post associated with a one-standard deviation shift to the right in respondent ideology. The bottom row shows the corresponding *differences* between engagement and attention gradients (Like – Read more, Share – Read more), capturing how much more strongly ideology predicts partisan-congruent likes and shares than partisan-congruent read-more clicks. Gray points represent study-specific estimates; pooled common-effect (fixed-effect) and random-effects (REML) meta-analytic estimates are shown by distinct point shapes (and colors). Error bars denote 95 percent confidence intervals.

share-read differences are about 0.49 and 0.53 (odds ratios ≈ 1.6 – 1.7), respectively, indicating that ideology predicts partisan congruence roughly 60–70 percent more strongly for likes and shares than for read-more clicks. In the representative U.S. sample, the corresponding differences are near 0.92 and 1.07 (odds ratios ≈ 2.5 and 2.9). The pooled meta-analytic differences range from about 0.55 to 0.63 for both likes and shares, implying odds ratios in the range 1.7–1.9. Across studies, ideological sorting is modest for attention but much stronger for active endorsement and sharing.

The smaller Δ estimates in the RIWI sample (Study 2) plausibly reflect differences in respondent composition and platform familiarity. For example, respondents reported fewer social-media accounts on average than U.K. Prolific respondents,

and attention-check pass rates were lower (58 percent in Study 2 vs. 85 percent and 91 percent in Studies 3 and 4).

When restricting to attention-check passers (Figure A7.1 in the online supplement), the substantive findings are the same. Full regression tables are also reported in the online supplement. As a robustness check, we expanded the pre-estimation exposure panel to include respondents who logged into the interface but never interacted with any tweet, assigning zeros for all outcomes. The ideology-congruence gradients are almost unchanged. This is consistent with the structure of the fixed-effects logit model: respondents with no outcome variation contribute little or no information to the within-person contrast between right- and left-aligned posts and may be dropped from the fitted model.

Finally, we re-estimate all our models across a range of simulated error regimes. These sensitivity analyses demonstrate that our results are robust to a mislabeling rate of ~40 percent in even the most anticonservative simulations. We describe these in detail in the online supplement.

Discussion and Conclusion

In this paper, we make two contributions. First, we show that, in a controlled social-media clone with a balanced political feed, the ideological diversity of what people *attend to* is substantially greater than what is implied by their likes and shares. When we measure attention as an intentional click to reveal additional content, information diets are markedly less ideologically congruent than when measured via outward engagement alone.

Second, we introduce a portable measurement strategy for attention. Our clone environment allows researchers to curate the content set, control feed ordering, and record both attention and engagement. We release the code and implementation details in the replication materials, contributing to recent efforts to move beyond coarse engagement proxies toward finer-grained measures of online information consumption (Green et al. 2025; Robertson et al. 2023).

Our approach carries trade-offs. The clone environment, while more naturalistic than standard surveys, is not identical to the platforms respondents ordinarily use; external validity therefore remains a central concern. The study standardizes the feed so that we can compare attention and engagement conditional on common exposure opportunities. Real platforms, by contrast, are shaped by prior decisions about whom to follow, which groups to join, and which sources to trust, as well as by ranking systems that may optimize for engagement, attention, revenue, or policy objectives in different ways.

A related concern is cue visibility. Although we classified posts using the truncated text shown to participants, the LLM and author coders may have inferred ideological slant more carefully than respondents did during a 120-second feed task. If the clone muffled ideological cues that are more obvious on X, Facebook, or other platforms—for example through source handles, images, replies, recommendation labels, or surrounding content—then the ideology-attention gradient on real platforms could be larger than we estimate here, and the gap between attention and engagement could be smaller. Conversely, if real feeds include many ambiguous or

low-salience political cues, or if political content is encountered incidentally among nonpolitical material, the gap could persist. The current design cannot adjudicate these platform-specific possibilities, and future work should vary cue obviousness, source labels, nonpolitical background content, and ranking rules directly. Our existing random and worst-case misclassification analyses address a related but distinct concern: whether the main engagement-attention contrast is fragile to errors in the left-right labels of posts.

Demand effects are also possible because every post in the clone is political, though existing evidence implies such effects are typically small in political-behavior settings (Mummolo and Peterson 2019). Measurement-wise, our attention metric is intentionally conservative: it does not record very brief glances at headlines or images that do not require expansion, and expansions may sometimes reflect curiosity about source or author as much as message content. We mitigate these issues by fixing the content set, randomizing content order and endorsement cues, and estimating exposure-level fixed-effects models that condition on tweet and person intercepts to isolate the ideology-by-content gradient.

Sampling is another limitation. Our studies rely on online panels (Prolific and RIWI) and a Prolific sample stratified to match key U.S. benchmarks. These are not probability samples of national populations. Nonetheless, recent work suggests Prolific yields strong data quality on comprehension, attention, and honesty (Eyal et al. 2021). The Prolific Representative sample further improves coverage on core demographics, but we should remain cautious about generalizing beyond the online, platform-using populations we study.

Substantively, our results are consistent with recent work showing that exposure and engagement can imply different levels of ideological segregation (Bakshy et al. 2015; González-Bailón et al. 2023; Nyhan et al. 2023; Robertson et al. 2023). We are more cautious, however, about generalizing from this controlled environment to any specific platform at a specific historical moment. The specific platform ecology once associated with Twitter may be changing, especially as X and other platforms alter moderation, verification, recommendation, and amplification systems. However, the broader phenomenon that motivates our design—out-of-network recommendation governed partly by attention signals—is not disappearing; if anything, it is increasingly central to contemporary social-media feeds, including short-video and recommendation-first platforms. Two implications follow. Estimates of ideological segregation that rely on engagement alone can overstate the insularity of information diets conditional on exposure; but real-world information diets also depend on upstream following decisions and algorithmic ranking systems that our design deliberately holds fixed. As platforms increasingly optimize for opaque combinations of attention and engagement signals, the divergence between what users read and what they endorse will remain difficult to infer from platform traces alone.

Future research can extend this design to causal questions by exogenously varying the ideological mix, source cues, or ranking rules to test how attention translates into downstream outcomes such as belief updating, outgroup affect, or online hostility. It will be useful to examine heterogeneity by political interest, partisanship strength, and media literacy, and to validate the attention measure against complementary indicators (e.g., dwell time and eye-tracking) where feasible.

As recommendation systems evolve, tracking the gap between attention and engagement over time can inform platform governance and audit efforts.

Notes

- 1 Appearances in the feed without either action constitute exposure only. These are sometimes referred to as “potential exposure,” “exposures,” or “impressions” (González-Bailón et al. 2023). We do not operationalize or record these.
- 2 In what follows, we variously refer to “tweets” or “posts.” We mean the same by both terms.
- 3 The rationale for this was that reference, in the present tense, to events that happened in the past may be confusing for participants. We instead aimed for a set of “evergreen” politically salient posts.
- 4 The landing screen is shown in Figure A3.1 in the online supplement.
- 5 Note that in the backend, each event is recorded individually as a click. This means a user might toggle the “like” or “share” buttons on and off, registering multiple engagements. In the main analyses, we filter out these events. When included, this has no substantive bearing on our findings.
- 6 We implemented this check by asking respondents in questionnaire whether they had seen tweets from individuals whose posts were not shown during the interaction window. Exact question wording and pass-rate computation are in the online supplement.
- 7 We acknowledge that it is unconventional to preregister a null hypothesis (though see Cortina and Folger 1998; Stanton 2021). Some argue it is not possible to confirm or reject such a hypothesis with null-hypothesis significance testing. In response, we focus our attention in the below on the substantive size of any attention association relative to the engagement gradient. Doing so allows us to gauge whether, in the presence of an association, it represents a meaningful size relative to a relevant baseline reference.
- 8 We initially amended our preregistration document to preregister more specific predictions about the absolute size of the attention and engagement gradients after the U.K. study, but later removed these amendments due to reader concerns about specifying hypotheses in terms of absolute magnitudes given plausible differences across the sample and data-generating processes across studies.
- 9 For more details of this sampling procedure, see <https://riwi.com/riwi-method-and-technology/>.
- 10 See <https://researcher-help.prolific.com/en/articles/445162-using-representative-samples-on-prolific-faqs>. We used the Prolific Political Representative sample, which “[Cross stratifies] on age (5 brackets), sex (2 groups) and political affiliation (3 groups) result[ing] in 30 subgroups: one for every combination of answers.”
- 11 We cannot verify that every respondent actually *viewed* every tweet. Tweets appearing earlier in the feed are plausibly more likely to be seen, but feed order is randomized at the respondent level. Under this design, any differences in which tweets are seen are expected to be unrelated to respondent ideology or tweet partisanship beyond random chance, so differential exposure should not systematically bias our estimates.
- 12 Some authors prefer the term “equal-effects” for the common-effect specification; see, for example, <https://wviechtb.github.io/metafor/reference/misc-models.html> in the discussion of the *metafor* package (Viechtbauer 2010).

¹³ Study 1 (United Kingdom, Prolific) began with 2,678 respondents; the baseline fixed-effects logit table corresponds to 2,257 respondents with modeled outcome variation. Study 2 (United States, RIWI) enrolled 856 respondents; the baseline fixed-effects logit table corresponds to 421 respondents with modeled outcome variation. Study 3 (United States, Prolific) enrolled 586 and 382 clicked at least once. Study 4 (United States, Prolific Representative) enrolled 352 and 236 clicked at least once. Figure A6.1 in the online supplement compares user-level behavior across studies.

References

- Bail, Christopher. 2021. *Breaking the Social Media Prism*. Princeton University Press. <https://doi.org/10.1515/9780691216508>
- Bail, Christopher A., Lisa P. Argyle, Taylor W. Brown, John P. Bumpus, Haohan Chen, M. B. Fallin Hunzaker, Jaemin Lee, Marcus Mann, Friedolin Merhout, and Alexander Volfovsky. 2018. "Exposure to Opposing Views on Social Media Can Increase Political Polarization." *Proceedings of the National Academy of Sciences* 115:9216–21. <https://doi.org/10.1073/pnas.1804840115>
- Bakshy, Eytan, Solomon Messing, and Lada A. Adamic. 2015. "Exposure to Ideologically Diverse News and Opinion on Facebook." *Science* 348:1130–32. <https://dx.doi.org/10.1126/science.aaa1160>
- Barberá, Pablo. 2015. "Birds of the Same Feather Tweet Together: Bayesian Ideal Point Estimation Using Twitter Data." *Political Analysis* 23:76–91. <https://doi.org/10.1093/pan/mpu011>
- Barberá, Pablo, John T. Jost, Jonathan Nagler, Joshua A. Tucker, and Richard Bonneau. 2015. "Tweeting from Left to Right: Is Online Political Communication More Than an Echo Chamber?" *Psychological Science* 26:1531–42. <https://doi.org/10.1177/0956797615594620>
- Barrie, Christopher. 2023. "Did the Musk Takeover Boost Contentious Actors on Twitter?" *Harvard Kennedy School Misinformation Review* 4. <https://doi.org/10.37016/mr-2020-122>
- Barrie, Christopher, Aybuke Atalay, and Alia ElKattan. 2025a. "An Enriched, Multimodal Social Media Dataset of a UK General Election Campaign." *Journal of Quantitative Description: Digital Media* 5. <https://dx.doi.org/10.51685/jqd.2025.017>
- Barrie, Christopher, Alexis Palmer, and Arthur Spirling. 2025b. "Replication for Language Models Problems, Principles, and Best Practice for Political Science."
- Bode, Leticia, Emily K. Vraga, and Sonya Troller-Renfree. 2017. "Skipping Politics: Measuring Avoidance of Political Content in Social Media." *Research & Politics* 4:2053168017702990. <https://doi.org/10.1177/2053168017702990>
- Cinelli, Matteo, Gianmarco De Francisci Morales, Alessandro Galeazzi, Walter Quattrociocchi, and Marco Starnini. 2021. "The Echo Chamber Effect on Social Media." *Proceedings of the National Academy of Sciences* 118:e2023301118. <https://doi.org/10.1073/pnas.2023301118>
- Cortina, Jose M. and Robert G. Folger. 1998. "When Is It Acceptable to Accept a Null Hypothesis: No Way, Jose?" *Organizational Research Methods* 1:334–350. <https://doi.org/10.1177/109442819813004>
- De Vreese, Claes H. and Peter Neijens. 2016. "Measuring Media Exposure in a Changing Communications Environment." *Communication Methods and Measures* 10:69–80. <https://doi.org/10.1080/19312458.2016.1150441>

- Downs, Anthony. 1957. "An Economic Theory of Political Action in a Democracy." *Journal of Political Economy* 65:135–150. <https://doi.org/10.1086/257897>
- Dubois, Elizabeth and Grant Blank. 2018. "The Echo Chamber Is Overstated: the Moderating Effect of Political Interest and Diverse Media." *Information, Communication & Society* 21:729–45. <https://doi.org/10.1080/1369118X.2018.1428656>
- Duchowski, Andrew T. 2007. "Visual Attention." Pp. 1–13 in *Eye Tracking Methodology: Theory and Practice*, Chapter 1. London: Springer.
- Eyal, Peer, Rothschild David, Gordon Andrew, Evernden Zak, and Damer Ekaterina. 2021. "Data Quality of Platforms and Panels for Online Behavioral Research." *Behavior Research Methods* 54:1643–62. <https://doi.org/10.3758/s13428-021-01694-3>
- Feldman, Lauren, Natalie Jomini Stroud, Bruce Bimber, and Magdalena Wojcieszak. 2013. "Assessing Selective Exposure in Experiments: The Implications of Different Methodological Choices." *Communication Methods and Measures* 7:172–94. <https://dx.doi.org/10.1080/19312458.2013.813923>
- Fishkin, James S. 1991. *Democracy and Deliberation: New Directions for Democratic Reform*. Yale University Press.
- Flaxman, Seth, Sharad Goel, and Justin M. Rao. 2016. "Filter Bubbles, Echo Chambers, and Online News Consumption." *Public Opinion Quarterly* 80:298–320. <https://doi.org/10.1093/poq/nfw006>
- Fletcher, Richard and Rasmus Kleis Nielsen. 2018. "Are People Incidentally Exposed to News on Social Media? A Comparative Analysis." *New Media & Society* 20:2450–68. <https://doi.org/10.1177/1461444817724170>
- Garrett, R. Kelly and Natalie Jomini Stroud. 2014. "Partisan Paths to Exposure Diversity: Differences in Pro-and Counterattitudinal News Consumption." *Journal of Communication* 64:680–701. <https://doi.org/10.1111/jcom.12105>
- González-Bailón, Sandra, David Lazer, Pablo Barberá, Meiqing Zhang, Hunt Allcott, Taylor Brown, Adriana Crespo-Tenorio, Deen Freelon, Matthew Gentzkow, Andrew M. Guess, et al. 2023. "Asymmetric Ideological Segregation in Exposure to Political News on Facebook." *Science* 381:392–8. <https://doi.org/10.1126/science.ade7138>
- Graf, Joseph and Sean Aday. 2008. "Selective Attention to Online Political Information." *Journal of Broadcasting & Electronic Media* 52:86–100. <https://dx.doi.org/10.1080/08838150701820874>
- Green, Jon, Stefan McCabe, Sarah Shugars, Hanyu Chwe, Luke Horgan, Shuyang Cao, and David Lazer. 2025. "Curation Bubbles." *American Political Science Review* 119:1704–22. <https://dx.doi.org/10.1017/S0003055424000984>
- Guess, Andrew, Kevin Munger, Jonathan Nagler, and Joshua Tucker. 2019. "How Accurate Are Survey Responses on Social Media and Politics?" *Political Communication* 36:241–58. <https://doi.org/10.1080/10584609.2018.1504840>
- Guess, Andrew M. 2015. "Measure for Measure: An Experimental Test of Online Political Media Exposure." *Political Analysis* 23:59–75. <https://doi.org/10.1093/pan/mpu010>
- Guess, Andrew M. 2021. "(Almost) Everything in Moderation: New Evidence on Americans' Online Media Diets." *American Journal of Political Science* 65:1007–22. <https://doi.org/10.1111/ajps.12589>
- Gutmann, Amy and Dennis F. Thompson. 2004. *Why Deliberative Democracy?* Princeton University Press. <https://doi.org/10.1515/9781400826339>

- Hickey, Daniel, Matheus Schmitz, Daniel Fessler, Paul E. Smaldino, Goran Muric, and Keith Burghardt. 2023. "Auditing Elon Musk's Impact on Hate Speech and Bots." Pp. 1133–37 in *Proceedings of the International AAAI Conference on Web and Social Media*. Vol. 17. <https://doi.org/10.1609/icwsm.v17i1.22222>
- Hobolt, Sara B., Katharina Lawall, and James Tilley. 2024. "The Polarizing Effect of Partisan Echo Chambers." *American Political Science Review* 118:1464–79. <https://doi.org/10.1017/S0003055423001211>
- Huszár, Ferenc, Sofia Ira Ktena, Conor O'Brien, Luca Belli, Andrew Schlaikjer, and Moritz Hardt. 2022. "Algorithmic Amplification of Politics on Twitter." *Proceedings of the National Academy of Sciences* 119:e2025334119. <https://doi.org/10.1073/pnas.2025334119>
- Lazer, David. 2020. "Studying Human Attention on the Internet." *Proceedings of the National Academy of Sciences* 117:21–22. <https://doi.org/10.1073/pnas.1919348117>
- Lazer, David, Eszter Hargittai, Deen Freelon, Sandra Gonzalez-Bailon, Kevin Munger, Katherine Ognyanova, and Jason Radford. 2021. "Meaningful Measures of Human Society in the Twenty-First Century." *Nature* 595:189–96. <https://doi.org/10.1038/s41586-021-03660-7>
- Lundberg, Ian, Rebecca Johnson, and Brandon M Stewart. 2021. "What Is Your Estimand? Defining the Target Quantity Connects Statistical Evidence to Theory." *American Sociological Review* 86:532–65. <https://doi.org/10.1177/00031224211004187>
- Messing, Solomon and Sean J. Westwood. 2014. "Selective Exposure in the Age of Social Media: Endorsements Trump Partisan Source Affiliation When Selecting News Online." *Communication Research* 41:1042–63. <https://doi.org/10.1177/0093650212466406>
- Metzger, Miriam J., Ethan H. Hartsell, and Andrew J. Flanagin. 2020. "Cognitive Dissonance or Credibility? A Comparison of Two Theoretical Explanations for Selective Exposure to Partisan News." *Communication Research* 47:3–28. <https://doi.org/10.1177/0093650215613136>
- Mummolo, Jonathan and Erik Peterson. 2019. "Demand Effects in Survey Experiments: An Empirical Assessment." *American Political Science Review* 113:517–29. <https://doi.org/10.1017/S0003055418000837>
- Narayanan, Arvind. 2023. "Understanding Social Media Recommendation Algorithms: Towards a Better Informed Debate on the Effects of Social Media." Knight First Amendment Institute at Columbia University. Accessed November 26, 2025 (<https://knightcolumbia.org/content/understanding-social-media-recommendation-algorithms>)
- Nyhan, Brendan, Jaime Settle, Emily Thorson, Magdalena Wojcieszak, Pablo Barberá, Annie Y. Chen, Hunt Allcott, Taylor Brown, Adriana Crespo-Tenorio, Drew Dimmery, et al. 2023. "Like-Minded Sources on Facebook Are Prevalent But Not Polarizing." *Nature* 620:137–144. <https://doi.org/10.1038/s41586-023-06297-w>
- Pariser, Eli. 2011. *The Filter Bubble: What the Internet is Hiding from You*. Penguin.
- Rettit API Contributors. 2023. "Rettit API: A Node.js Library for Interacting with Twitter's API." Version 2.0.1 (<https://www.npmjs.com/package/rettit-api/v/2.0.1>).
- Robertson, Ronald E., Jon Green, Damian J. Ruck, Katherine Ognyanova, Christo Wilson, and David Lazer. 2023. "Users Choose to Engage with More Partisan News Than They Are Exposed to on Google Search." *Nature* 618:342–48. <https://doi.org/10.1038/s41586-023-06078-5>
- Stanton, Jeffrey M. 2021. "Evaluating Equivalence and Confirming the Null in the Organizational Sciences." *Organizational Research Methods* 24:491–512. <https://doi.org/10.1177/1094428120921934>

- Stroud, Natalie Jomini. 2010. "Polarization and Partisan Selective Exposure." *Journal of Communication* 60:556–76. <https://doi.org/10.1111/j.1460-2466.2010.01497.x>
- Stroud, Natalie Jomini. 2017. "Attention As a Valuable Resource." *Political Communication* 34:479–89. <https://doi.org/10.1080/10584609.2017.1330077>
- Süllflow, Michael, Svenja Schäfer, and Stephan Winter. 2019. "Selective Attention in the News Feed: An Eye-Tracking Study on the Perception and Selection of Political News Posts on Facebook." *New Media & Society* 21:168–90. <https://doi.org/10.1177/1461444818791520>
- Team, PostHog. 2024. "posthog-js: JavaScript Library for PostHog Analytics." Accessed July 26, 2024 (<https://github.com/PostHog/posthog-js>).
- Viechtbauer, Wolfgang. 2010. "Conducting Meta-Analyses in R with the metafor Package." *Journal of Statistical Software* 36:1–48. <https://doi.org/10.18637/jss.v036.i03>
- Ye, Jinyi, Luca Luceri, and Emilio Ferrara. 2024. "Auditing Political Exposure Bias: Algorithmic Amplification on Twitter/X Approaching the 2024 U.S. Presidential Election." <https://arxiv.org/abs/2411.01852>.

Acknowledgments: We thank audiences at the NYU Center for Social Media and Politics for their feedback on earlier versions of this paper. This research was supported by a British Academy/Leverhulme Small Research Grant (SRG2223 230865). OpenAI GPT-4o and GPT-5 and Anthropic Claude (Sonnet 4.5) were used for code, code review, and manuscript redrafting. The authors declare no competing interests. Corresponding author: Christopher Barrie (cb5691@nyu.edu).

Author Contributions. Conceptualization: C.B., A.A. Methodology: C.B., A.A., A.E. Data curation: C.B. Data visualization: C.B. Writing—original draft: C.B., A.A., A.E. All authors approved the final submitted draft.

Christopher Barrie: Department of Sociology, New York University.
E-mail: cb5691@nyu.edu

Aybuke Atalay: Department of Politics and International Relations, University of Edinburgh. E-mail: aybuke.atalay@ed.ac.uk

Alia ElKattan: Department of Politics, New York University. E-mail: alia.elkattan@nyu.edu