

Digital News Consumption: Evidence from Smartphone

Content in the 2024 U.S. Elections

Guy Aridor, Tevel Dekel, Rafael Jim'enez-Dur'an, Ro'ee Levy, Lena Song

Working paper 11-2026

The Foerder Institute for Economic Research
and

The Sackler Institute of Economic Studies

Digital News Consumption: Evidence from Smartphone Content in the 2024 U.S. Elections

Guy Aridor^{1,*}, Tevel Dekel², Rafael Jiménez-Durán³, Ro’ee Levy^{2,4}, Lena Song⁵

¹Northwestern University, Kellogg School of Management. ²Tel Aviv University. ³Bocconi University, IGIER, and Stigler Center. ⁴CEPR. ⁵University of Illinois Urbana–Champaign.

Smartphones are a central gateway to news, yet we know little about the political content people see on their phones. We use smartphone data to measure election-related exposure across applications during the 2024 U.S. presidential campaign. We document three main findings. First, exposure to election-related content was surprisingly low. On an average day, the median individual in our data was exposed to substantially less election-related content than appears in a single New York Times article. Second, exposure was largely decoupled from traditional news consumption. News apps accounted for less than 1% of exposure, while social media accounted for over thirty times as much. Third, exposure was highly unequal: Individuals in the 90th percentile were exposed to more than 50 times as much election-related content as those in the 10th percentile. A variance decomposition shows that the applications people use help explain the variation in exposure, but individual differences remain the dominant driver.

Democratic accountability requires voters to encounter information about candidates and issues during election campaigns, often through exposure to political information in the media [1, 2, 3]. Smartphones have recently become the primary method people first come across news [4]. Still, we know surprisingly little about the political content people actually see on their phones.

The implications of smartphone-based news exposure are ambiguous. In terms of quantity, it is seemingly easier than ever for voters to find information, and even passive voters may encounter election-related information through social media sharing or targeted advertising.¹ But smartphones also provide abundant non-news alternatives, which may crowd out political attention [8]. In terms of sources and quality, smartphones may shift exposure away from editor-curated news outlets toward non-traditional channels, particularly social media, where algorithms may surface more polarized, toxic, or less reliable content [9, 10, 11]. Finally, these algorithms may also matter for exposure inequality. They may provide users with content matching their preferences and thus amplify differences between news seekers and news avoiders [12]. Alternatively, they may reflect platform priorities (as illustrated by Meta’s reported downranking of political content and X’s apparent movement in the opposite direction) and create differences in exposure across platforms [13, 14].

These competing possibilities raise three empirical questions regarding election-related exposure on smartphones. First, how much content are people exposed to? Second, where does this exposure occur? Third, what drives differences in exposure across people? Existing evidence cannot answer these questions directly because surveys are vulnerable to substantial measurement errors [15, 16] and browser data miss most in-app smartphone consumption. For example, surveys suggest that social media plays an important role in news consumption [4], while browsing data suggest that social media accounts for only a small share of visits to news sites [17, 18]. Neither approach directly observes the political content people see inside social media feeds. Single-platform studies help overcome this challenge by observing the feed of one platform [19], but they cannot compare exposure across the many applications people use on their phones.

We address these questions using data that captures keyword occurrences across all applications on thousands of Americans’ smartphones during the consequential 2024 U.S. presidential election. Our data covers over 500 pre-specified election-related keywords, including the names of candidates. We observe the exact time, app, and device where each keyword was encountered, allowing us to study both the extent of exposure throughout the campaign and the sources of exposure. This approach builds on the Screenomics methodology to measure digital behavior [20, 21, 22] while scaling and extending the method to study an election campaign.

Our first finding is that election-related smartphone exposure was strikingly low. For example,

¹Campaigns spent around two billion dollars on digital ads in the 2024 election, and more candidates advertise on social media, which is mostly used on smartphones, than on television [5, 6, 7].

the median individual in our data was exposed to substantially less election-related content than they would have encountered by reading a single election-related *New York Times* article. This suggests that (the lack of) media exposure may be one mechanism contributing to limited and unequal political knowledge across individuals and issues [23, 24].

Our second main finding is that exposure rarely came from traditional news channels and instead often occurred through non-traditional unvetted sources, especially social media. This implies that the common practice of inferring exposure from the sources (e.g., domains) that people visit misses most of the political content people are exposed to on their phones [25, 26, 27]. Both findings depart from expert perceptions. Drawing on a survey of political economists, political scientists, and experienced forecasters, we find that these experts substantially overestimated total exposure to election-related content and the share of exposure occurring on news apps.

Finally, we find substantial variation across individuals and applications in exposure to election-related content. Using a variation decomposition exercise and an original survey, we study whether this variation is due to systematic differences across individuals (driven by algorithmic personalization or differences in individuals’ usage patterns) or systematic differences across applications (driven by algorithmic prioritization or differentiation in content creation). Our analysis contributes to the debate on whether algorithms shape political exposure, which has primarily focused on single-platform settings [28, 29, 19, 30]. Our results suggest that algorithms play a role in driving exposure, but that interventions focusing on algorithms alone may have limited effects if they do not also address individual demand for political content.

1 Measuring smartphone exposure to election-related content

Our primary data consists of observed keyword occurrences across all apps on individuals’ smartphones for several thousand Americans. Most of our analysis focuses on the election campaign from September 1st to November 5th (election day).

The data comes from Screenlake, a company that creates and maintains a proprietary database of term encounters occurring in natural smartphone usage, primarily to serve enterprise businesses in measuring their brands’ organic popularity. Screenlake’s data is sourced from an SDK (software development kit) that is embedded within a popular smartphone application on Android.² With users’ permission, the SDK passively checks for the presence of specific keywords every three seconds continuously across apps. For each keyword encountered, we observe the device, the exact time when the keyword appeared, and the app that was on-screen at the time.³ In addition to

²Android is used by approximately 40% of Americans. Similar functionality is not available to third-party developers on iOS.

³All such content processing is done on the smartphone device itself to avoid the collection of personally identifiable or sensitive information.

observing text rendered visible on-screen, we observe multimedia: keywords that arise in the ‘alt text’ that briefly describes a video or image, keywords that appear in video captions, and for some applications, keywords that appear in the transcriptions of videos.⁴

We observe two main groups of election-related terms: 1) *election terms*, which are words related to the elections (“polls”, “ballot”, “campaign”, “debate”, “democracy”, “Republican”, “Democrat”, “Congress”, “swing”, “vote”, “election”, “electoral”, “Mail-In”) and 2) *politicians*, which includes current and former presidential and vice-presidential candidates (“Joe Biden”, “Robert F Kennedy Jr”, “Nikki Haley”, “Kamala Harris”, “Donald Trump”, “Tim Walz”, “JD Vance”, “Hillary Clinton”, “Bernard Sanders”), and virtually all congressmembers and governors.⁵ Supplementary Section S9.1 provides the full list of over 500 distinct election-related keywords provided to us by Screenlake. In addition to the election-related keywords, we observe exposure to several thousand control keywords covering sports, entertainment, commercial brands, and other topics, which we use as benchmarks and to proxy total time spent on each app. We manually classify apps into the following categories: communication, social media, music and video, news, search, browser, and AI assistant, as described in Supplementary Section S9.2.

Primary Outcomes We measure consumption using two main measures. First, we measure election-related *encounters*, defined by the number of occurrences of an election-related keyword that a user encountered. Second, we measure election-related *exposure*, defined as the number of observed time periods with at least one keyword occurrence of our election-related keywords. This is less sensitive to the exact combination of election-related keywords used as long as at least one relevant keyword is observed. It also allows us to interpret our results in units of time. As keywords are generally captured every three seconds, we make a conservative assumption that the keyword was observed for all three seconds, and we multiply the number of exposures by three to measure exposure in seconds.

Keyword Informativeness Our measures of consumption rely on the assumption that the pre-specified keywords are informative election-related content. In Supplementary Section S1 we formulate the measurement problem as a classification problem where the presence of at least one of our keywords (an “exposure”) indicates that an individual is consuming election-related content at a given observation period. We primarily rely on two supply-side benchmarks – the set of articles from the New York Times and Fox News during the election period – to benchmark

⁴Alt text is often supplied to the app by humans or AI as part of a general push for accessibility. It is common, but not always available for all images and videos. Transcriptions are common for videos on YouTube, but not, for instance, podcasts on Spotify.

⁵For political figures with names that are not especially distinct in American English, Screenlake maintains proprietary logic for ascertaining if a reference is indeed in reference to the intended political figure.

the performance of our keyword-based classifier. These supply-side benchmarks provide us with ground-truth labels for whether content is election-related since they provide article-related tags, which allows us to differentiate between articles that are about the election versus other topics. Using these ground-truth labels we then measure the performance of our keyword-based classifier at correctly detecting election-related content.

We demonstrate that our keywords perform reasonably well, achieving high recall with 97.74% and 97.89% of election-related articles on the New York Times and Fox News, respectively, including at least one of our keywords.⁶ Furthermore, we demonstrate that the keyword-based classifier has a moderately high precision statistic of 0.58 and 0.70 for the New York Times and Fox News, respectively. Combined, these statistics suggest that our keyword-based classifier can accurately measure the extent of election-related content consumption.⁷ If anything, the moderate precision suggests that our results may be an upper-bound on the magnitude of consumption.⁸

The Panel Our panel consists of active users who installed a specific app. An important concern is how findings from this group can be extrapolated to the broader U.S. population. We conduct several exercises, including collecting app time use from a representative sample during the same time period, to assess the generalizability of our results. We document the details of these exercises in Supplementary Section S2 and provide here a brief summary.

First, we reweighted our sample to match the U.S. population based on inferred gender and age and find that our results stay qualitatively the same. We also find that reweighting our sample based on the inferred state does not substantially change the results. Our sample may still differ from the population in unobserved ways, such as political interest. Therefore, our second exercise compares the app composition and phone usage of our sample to market-level Google Play application downloads data and industry benchmarks. We show that our sample is similar to the population at the extensive margin, including the installation of apps related to news. Third, we compare our sample’s app and total phone usage at the intensive margin to a representative sample of Prolific

⁶We additionally compute this statistic using 250-word increments – an approximation for how many words would be on an individual’s screen at a single point in time. Using these 250-word increments, we still find a high recall of 0.87 and 0.88 on the New York Times and Fox News, respectively. In other words, if someone has just one random segment from a New York Times or Fox News article on their smartphone screen, we are highly likely to record an exposure to election-related content.

⁷These results demonstrate that our keywords are highly informative about classifying election-related content written by journalists and curated by editors. We additionally compute recall on a benchmark dataset of user-generated content by measuring its performance over the full set of comments on the sub-reddit ‘r/politics’. We find that 90% of the posts have at least one of our keywords in their comment thread.

⁸Selecting a keyword subset from our available bank corresponds to choosing a point on the precision-recall curve. While, in principle, we could prioritize a smaller set of keywords with higher precision and lower recall, in our context the cost of false negatives is higher than that of false positives. We therefore selected a set that lies toward the high-recall region of the curve, accepting more false positives in exchange for fewer cases where we miss a relevant keyword.

participants we recruited in the week before election day. We find that the average user in our data has similar phone usage, spends more time on social media applications, and, if anything, is more likely to have news apps installed than the average American. Fourth, we study the primary reason users install the application—managing their screen time. We conducted an additional survey and found no evidence that users of screen time management software are less politically engaged; in fact, individuals who actively manage their screen time may be slightly more politically engaged. All these patterns suggest that our finding that election-related consumption is low does not stem from low phone usage or lack of news interest, compared to the US population.

2 How much election-related content are people exposed to?

2.1 Election-related consumption is low

Figure 1 presents the daily consumption of content during the election campaign for the median individual. On an average day between September 1st, 2024 and November 4th, 2024 (the day before election day), the median individual only had 13.2 encounters with our set of election-related keywords, as shown by Panel A. This is the total number of times these election-related words were mentioned in any app, including, for example, social media posts, news articles, or someone mentioning the word ‘Trump’ in a text message. To benchmark this result, we compare it to the universe of New York Times articles described in Section S9.1 and find that these encounters correspond to 57.2% of the terms that appear in an average political article. This means that someone reading a single political article would encounter almost twice as many keywords as the median individual on an average day even without seeing any political ads, political posts on social media, or discussing politics on their phone.

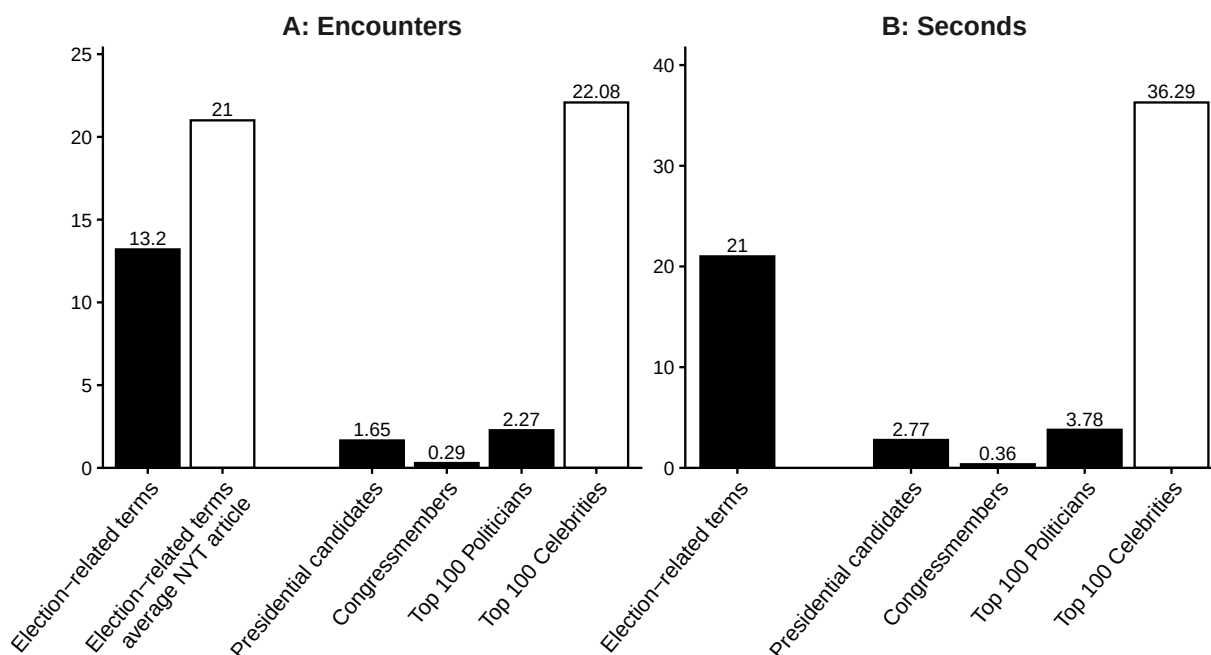
To further interpret the meaning of these encounters, we distinguish between incidental, isolated encounters and sustained encounters based on whether they occur near other election-related encounters within a ten-minute window. Encounters that appear alone, or with only one other encounter, are more likely to reflect incidental exposure, such as seeing a social media ad or a recommended video. By contrast, encounters that appear in larger clusters are more likely to reflect sustained engagement with election-related content, such as watching a political video, participating in a conversation, or reading an article. Using this classification, we find that for the median individual, around 60% of election-related encounters are incidental, suggesting that most encounters are unlikely to reflect sustained interest in the election.

Using our time-based measure, the right panel of Figure 1 shows that the median individual was exposed to approximately 21 seconds of election-related terms on their phone. The median individual spends on average approximately 5.54 hours a day on their phone, suggesting that indi-

viduals saw election-related terms only 0.1% of the time.

In Supplementary Section S3.1 and Appendix Table 1, we show that our results are robust to reweighting the sample based on demographics to match the US population, excluding Spanish speakers, expanding keyword definitions to include political issues, accounting directly for the precision and recall of our classifier, and making alternative measurement assumptions.

Figure 1: **Election-related content consumption is low.**



NOTES: The figure reports the median encounters (Panel A) and seconds of exposure (Panel B) among active devices on an average day between September 1 and November 4, 2024. From left to right, each panel presents the overall consumption based on the full set of election-related terms, the number of election-related terms in an average New York Times article (only for encounters), names of presidential candidates, congressmembers, top 100 politicians, and top 100 celebrities.

2.2 Exposure to congressmembers is even smaller

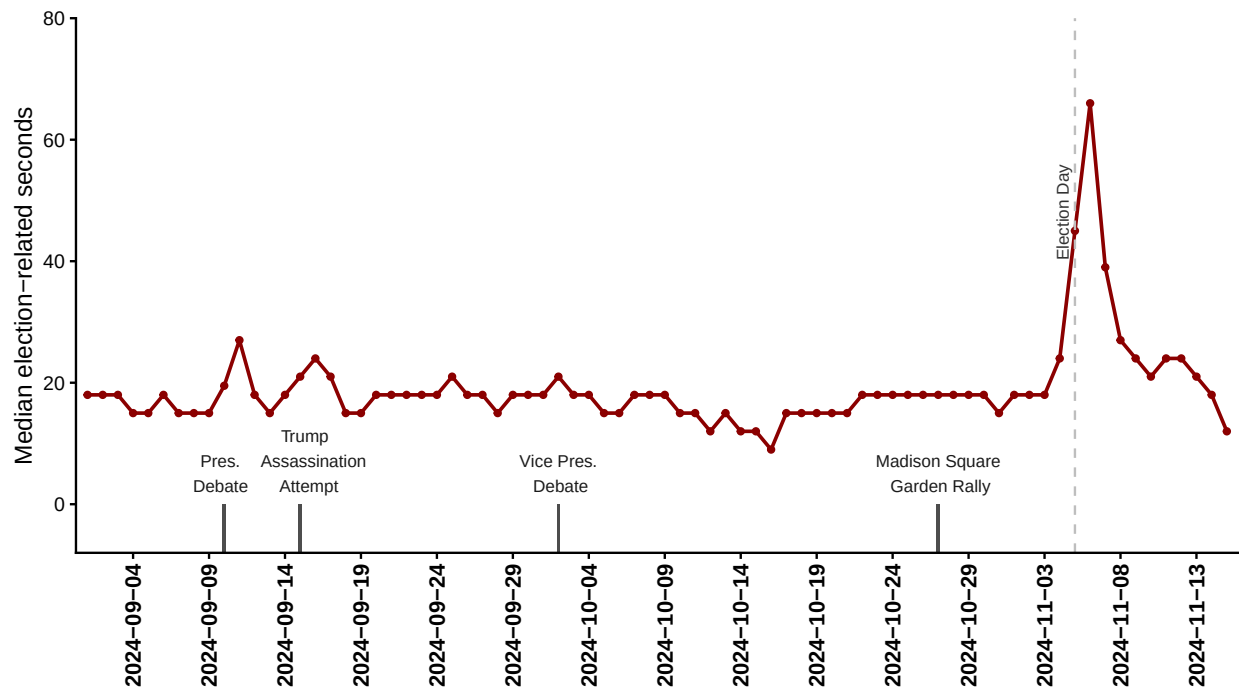
The results so far are based on both election terms and politicians' names. Figure 1 shows that the median individual was exposed to the president and presidential candidates ('Donald Trump,' 'Joe Biden,' and 'Kamala Harris') for less than 3 seconds on an average day. Exposure to the top 100 politicians in our keyword list was slightly higher but still only 3.8 seconds. As a benchmark, we compare politicians to celebrities and find that exposure to the 100 most observed celebrities was 9.6 times greater than exposure to the 100 most observed politicians. This finding is unlikely to suffer from measurement error. First, while people may use different terms to discuss the election, the names of politicians are typically well-defined and should be easily captured. Second, the

comparison to celebrities assuages concerns that we were not able to capture keywords from users’ phones due to some technical limitation, since the names of celebrities were captured using the exact same method.

While exposure to presidential candidates was limited, exposure is minuscule when focusing on congressmembers. Surveys have shown that individuals know little about their congressmembers [31, 32]. Therefore, exposure to information about congressmembers is especially important as voters probably have weak priors, and informing them about their representatives may affect their decisions. However, Panel B in Figure 1 shows that the median individual was exposed for only 0.4 seconds per day to *all* congressmembers (in 2024, a serving congressmember ran in approximately 89% of districts). Indeed, this number is so low that on an average day only 18% of individuals encountered any congressmembers, as shown in Supplementary Table S2.

2.3 Exposure is remarkably stable throughout the campaign

Figure 2: Election-related exposure is remarkably stable over the campaign.



NOTES: The figure displays the median daily number of seconds of exposure to election-related terms per active user between September 1 and November 15, 2024. Key campaign events are indicated along the x-axis: the Presidential Debate between Kamala Harris and Donald Trump (September 11), the second Trump assassination attempt (September 15), the Vice Presidential Debate between Tim Walz and J. D. Vance (October 2), and Trump’s Madison Square Garden rally (October 27). The vertical dashed line marks election day (November 5).

Given the low consumption of election-related content, a natural question is whether the level of consumption gradually rose closer to the election, or spiked after the many notable events in the campaign, such as the second attempt to assassinate Trump, the presidential and vice presidential debates, and Trump’s Madison Square Garden rally. Figure 2 plots the median election-related exposures in seconds over time and shows that it is remarkably stable (Supplementary Fig. S6 shows similar results using encounters). In 80% of observed days, the median user’s daily exposure to political content ranged between 15 and 20 seconds. This confirms that exposure was low throughout the campaign and not only at the beginning of the time period we analyze.

There are two exceptions to this trend: the presidential debate and election day. On the day of the first presidential debate between Kamala Harris and Donald Trump, exposure increased by 44% compared to the previous week. The increase in consumption due to the debate dovetails previous literature showing that debates draw large audiences [33]. On election day, there was a dramatic increase where the median exposure time increased by 287% from the average in September-October to the peak occurring the day after the election. In Appendix Figure 6 we zoom in on election day and show that most election-related exposures occurred *after* polls closed. This result indicates that most of the increase in exposures around election day is not driven by individuals seeking information before voting, but rather people reading about the results. To summarize, with the sole exception of the presidential debate, there were no large population-level shocks that induced additional information acquisition before voting (though events can differentially impact specific subgroups; see Supplementary Section S6 on Taylor Swift’s endorsement).

3 Where does election-related exposure occur?

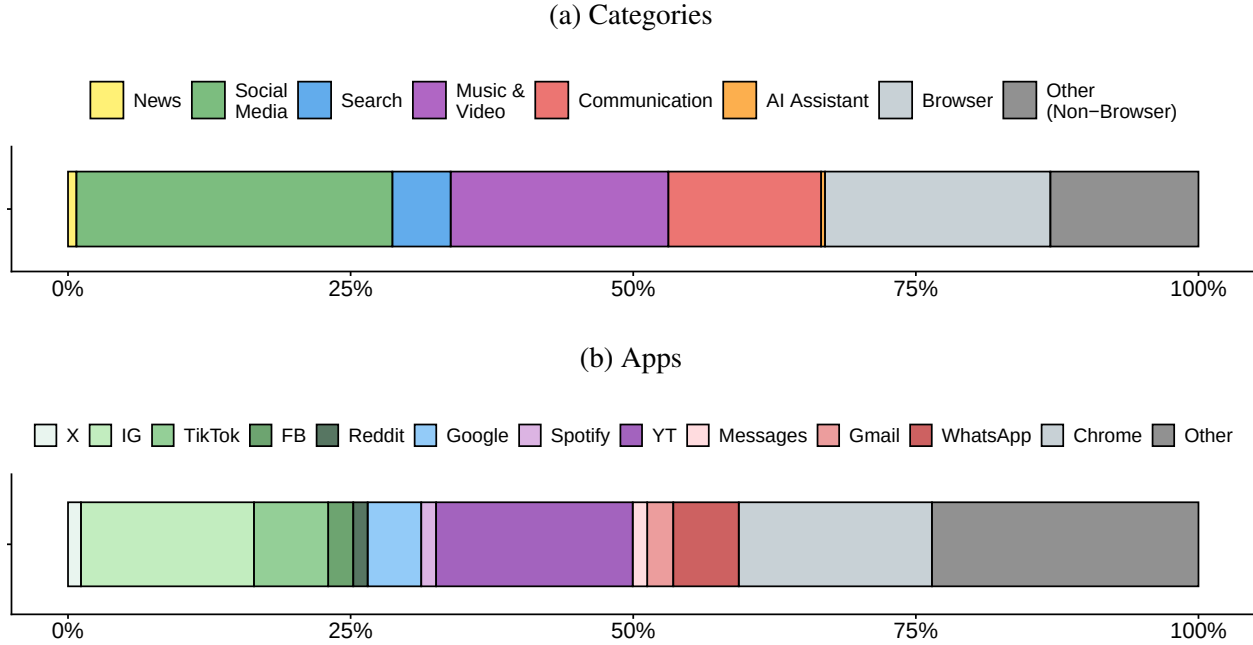
3.1 The vast majority of consumption occurs outside of news applications

There is concern that individuals today receive less political information from editor-curated news sources and instead encounter content in non-traditional channels, such as algorithmically curated social media feeds and music & video apps [34, 35]. A shift towards non-traditional sources of information could affect the *quality* of information, as these environments may surface more polarized, toxic, or less reliable material [11, 36, 10]. However, studies using web browsing data did not find evidence for the hypothesis that social media has become a prominent source of news exposure [17, 37]. Our screen-level measure of exposure across all smartphone apps (including typically private channels such as communication) provides a clear test of this hypothesis.

Figure 3a shows that individuals rarely consume election-related content through news applications. The mean individual-level share of exposures from news applications is only 0.74%. We find that for every second of exposure in a news app, there were approximately 35 seconds of

exposure on social media apps.

Figure 3: **Election-related exposure occurs predominantly outside news apps.**



NOTES: The figure shows the overall distribution of election-related exposures across apps and app categories during the period September 1–November 4, 2024. Panel (a) displays the share of exposures by category, and Panel (b) by app. Each bar shows the average share of a user’s total election-related exposures occurring in that category or app. The shares are computed in two steps: first, for each user, exposures are summed by category/app and divided by the user’s total election-related exposures; second, these user-level shares are averaged across all devices. Abbreviations refer to Instagram (IG), Facebook (FB), and YouTube (YT).

One limitation is that we cannot observe the websites individuals visit through their mobile browsers, and these websites can belong to any category (e.g., facebook.com or nytimes.com). In Supplementary Section S3.3.1, we describe two methods for imputing news-site consumption on the browser, and find that the mean share of exposures through news apps and visits to their websites on browsers is still only 0.89%-4.93%.⁹

In Supplementary Table S2, we report consumption by categories at the extensive margin and find that most people are *never* exposed to election-related content on news applications. Throughout the entire campaign period, only 9% of individuals encountered election-related content on news apps.

Where do people get election-related content? Exposure via communication and search is

⁹Individuals may still encounter content from news outlets on social media. However, these encounters are quite different from using news apps. Individuals typically observe only the headline, the content is algorithmically curated, and tends to be more like-minded and extreme [11].

important as these channels are more likely to capture active interest. Figure 3a shows that 18.67% of exposure is from communication and search. While this is much higher than news apps, it is relatively low compared to the social media and music & video, which make up 47.20% of total exposure. This result reinforces our conclusion that a substantial share of news is consumed passively. Indeed, as Supplementary Table S2 shows, on an average day only about one-third of individuals observe even a single election-related term in their communication apps.

3.2 Experts overestimate both exposure and news-app use

To calibrate our findings against expert judgment, we conducted a survey of researchers with expertise in political economy, political behavior, and related fields, as well as experienced forecasters, following [38]. Survey details and results from 223 respondents are reported in Supplementary Section S8.

Appendix Figure 7 shows that our results are far from obvious—even experts systematically overestimate how much political content people actually consume and underestimate the relative importance of entertainment and non-traditional news sources. While actual median exposure to election-related content is less than one minute, the median expert prediction is 45 minutes. In addition, experts underestimate the relative exposure to celebrities compared to politicians. While the median individual was exposed to 9.6 times more celebrities than politicians, the median expert prediction for this ratio is 3.5, with only 8.1% of experts predicting that the ratio was between 8 and 12. Experts are more accurate when asked about the composition of election-related exposures as most predict that a minority of exposure comes from news sites or apps. Still, experts overestimate the extent of consumption from news sites or apps, with a median guess of 15%, while we find that the actual share is less than 5%.

4 What drives the differences in exposure to the elections?

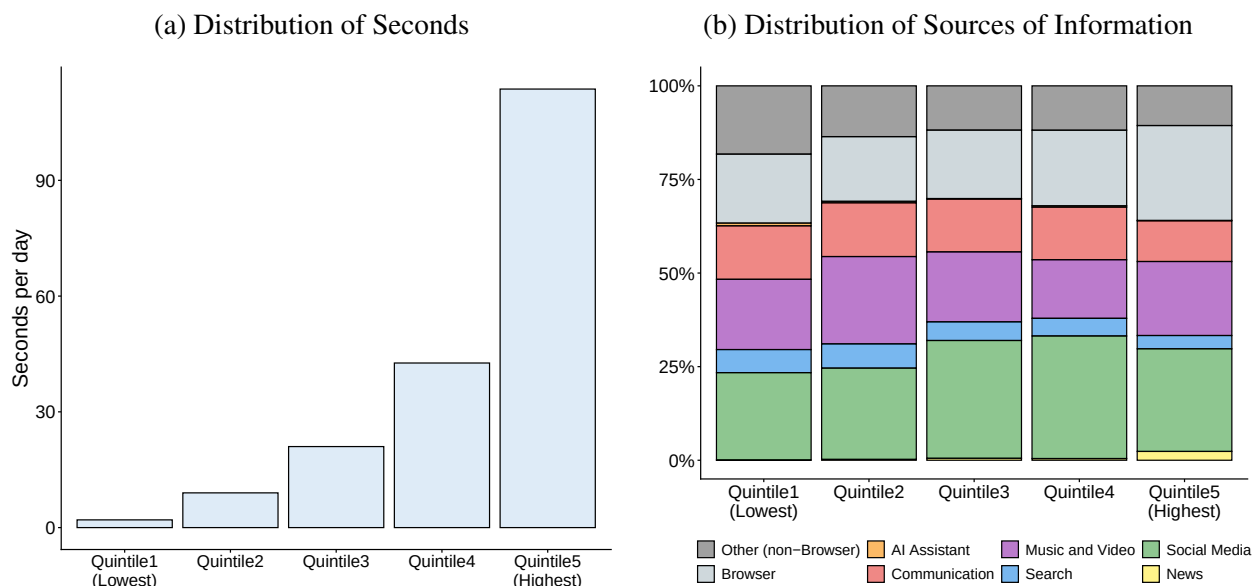
While the median individual consumes a small amount of election-related content, there are intrinsic differences across individuals in their interest in such content and the applications they use. As such, we next turn to characterizing the extent of heterogeneity in and quantifying the potential drivers of these differences in consumption.

4.1 Exposure varies sharply across individuals and applications

Figure 4a shows the average daily exposure to election-related content across quintiles, revealing substantial heterogeneity, mirroring previous findings on misinformation exposure [39]. The second quintile averages 9 seconds of daily exposure, compared to 43 seconds in the fourth quintile.

The gap widens further at the extremes: individuals at the 10th and 90th percentiles have around 2 and 114 seconds of exposure, respectively, corresponding to a 90/10 percentile ratio of 56.8. Appendix Figure 8 plots the Lorenz Curve of election-related consumption and highlights that the inequality in exposure is dramatic, with a higher Gini coefficient than phone usage and TV watching. This Gini coefficient is also higher than the one for US income [40] and approximately the same as the one for world income [41]. Journalists, pundits, and politicians are probably overrepresented among the heavy news consumers and thus may have a skewed view of the electorate’s news consumption, providing one potential explanation for the incorrect beliefs of the experts we surveyed.

Figure 4: **Heterogeneity in election-related content consumption is dramatic.**



NOTES: Panel (a) shows the distribution of average daily exposure in seconds to election-related content across device quintiles, ranked by overall exposure levels. Exposure for each device is computed as total election-related exposure divided by the number of active days between September 1 and November 4, 2024. Devices are then sorted by this average exposure and divided into five equal-sized quintiles. Panel (b) shows the composition of these exposures by source category (*e.g.*, social media, news, search, communication, browser, and other apps). Category shares are computed at the user level as the fraction of that user’s total election-related exposure attributable to each category and then averaged within quintiles. Quintiles are defined as in Panel (a), based on a user’s average daily exposure seconds.

Importantly, exposure varies based on individual characteristics. Because the U.S. relies on the Electoral College for presidential elections, swing states play a pivotal role in determining outcomes. Individuals in swing states likely have greater incentives to seek election-related information or are more heavily targeted by political messaging. Appendix Figure 9a shows that the median individual in a swing state receives 88% more election-related exposures than other states.

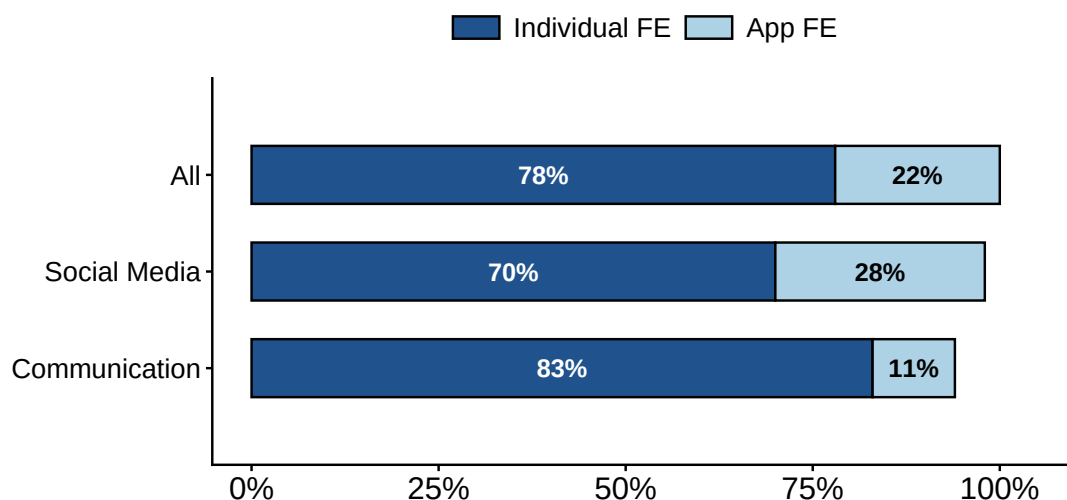
As both the quantity and quality of exposure matter, Figure 4b examines how the composition

of election-related content varies across the exposure distribution. While non-traditional sources, which may provide lower-quality content, are the dominant source of information across all groups, individuals with lower overall exposure rely primarily on these sources and consume virtually no election-related content from news applications. Moreover, Appendix Figure 9a shows that individuals using news apps are exposed to more election-related content. These patterns suggest that inequality in the quality of election-related exposure may mirror inequality in quantity.

After reviewing differences across individuals, in Appendix Figure 9b we compare the share of election-related content on different social media applications. We document substantial heterogeneity, with applications such as X and Reddit having considerably higher shares of election-related exposures than Facebook and TikTok.

4.2 Individual differences drive most of the variation in exposure, but social media applications matter too

Figure 5: **Individual effects dominate, but application effects are non-trivial within social media.**



NOTES: The figure shows the share of explained variance in the share of election-related exposures across all apps (excluding those without categorization), social media apps, and communication apps. We decompose the relative importance of the terms in (S2) and report the individual and app fixed effects as they are the dominating terms. The values for each of the terms are reported in Appendix Table 2, showing most are negligible apart from the individual and app terms, and summing to 100% when including all of the terms. The adjusted R^2 values are 0.26 (All), 0.27 (Social Media), and 0.21 (Communication). The number of apps included in each category is 285, 10, and 13, respectively, based on a total of 1,469 individuals.

We assess the relative importance of two dimensions – individuals and applications – as drivers of the observed inequality in content consumption. Decomposing these forces is important to charac-

terize what policy interventions can increase or equalize election-related content consumption. If differences across applications are the main driver of variation in exposure, policy responses could include algorithmic audits or encouraging individuals to diversify the apps they use [42, 43]. If, instead, individual heterogeneity dominates, understanding what motivates political information acquisition becomes central (e.g., changes to the format of news, its price, its perceived relative trustworthiness, or the stakes of being informed [44, 45, 46, 47]).

To disentangle these forces, we leverage the individual-level variation in election-related exposure across applications, and conduct a variance decomposition in the spirit of [48]. Supplementary Section S4 provides more details.

Figure 5 shows that individual-specific heterogeneity accounts for 3 times more of the explained variance than app-specific heterogeneity when aggregating across apps. This figure also shows that the relative role of apps is higher in the case of social media—where factors such as algorithmic biases can allow for app-specific systematic effects across individuals. In contrast, as expected, the role of apps is smaller for communication apps where individuals actively decide what to discuss and content curation is less pervasive. Appendix Table 2 contains the full decomposition results and Supplementary Section S5.3 shows that the results are robust across alternative specifications, including the bias correction of [49].

4.3 Social media apps can influence exposure through content prioritization

The role of social media applications in driving exposure raises a natural question: do these app-level effects reflect users employing different platforms in different ways, or do platform algorithms systematically prioritize different content? We present three pieces of evidence indicating that platform prioritization plays a substantive role.

First, consistent with news reports on the differences between how X and Meta prioritized election-related content during the campaign [13, 14], Appendix Figure 10 shows that X users are consistently over-exposed to election-related content compared to other applications, while Facebook users are under-exposed. These estimates take into account individual fixed effects and thus are not affected by differences in the apps’ users. Second, we examine a keyword that is likely to only reflect app prioritization and not differences in user preferences: ‘Elon Musk’. We find that this keyword appears more often on X than on any other social media app, consistent with media reports of Musk prioritizing his posts on X [50]. The magnitude is substantial, with Elon Musk being 20 times more likely to appear on X compared to Facebook.

Third, to more cleanly distinguish between the two candidate explanations, we conducted an additional survey measuring whether participants were exposed to too much or too little content across topics and applications, allowing us to compare exposure to people’s preferences. Partic-

ipants report greater *over*-exposure to political content and to Elon Musk on X relative to other applications, suggesting that these terms appear often due to the prioritization of these topics by the platform and not because of user preferences. More details are provided in Supplementary Section [S5.4](#).

5 Discussion

In this paper, we document new facts about news consumption during a highly consequential election season using novel data on smartphone content. We find that overall exposure remains limited and stable in the two months leading up to the election. The little exposure that does occur is largely centered on national politics, and a significant portion of it comes from social media and video apps. While the median individual consumes very little, we find substantial heterogeneity in consumption, with individual differences, rather than app-level differences, emerging as the dominant driver of variation in consumption.

These findings have important implications for policy. First, the low quantity and quality of election-related content consumption indicate a potential mechanism for why existing research has found low levels of political knowledge. Both the level of consumption and the share coming from news sources fall well short of what experts consider ideal. The median expert in our survey believes people should spend around 30 minutes per day consuming election-related content, and nearly all (93.8%) say that news apps should be the primary source for making an informed decision about the election.

Second, to the extent that policymakers wish to increase election-related content consumption, our variance decomposition results suggest that it may be more effective if they focus on changing individual incentives to consume news, rather than primarily regulating algorithms or content prioritization practices. Nonetheless, our results show that application-level prioritization (particularly on X) can meaningfully shift exposure, suggesting caution about the potential agenda-setting power of these applications and underscoring the need for cross-platform measurement in research on social media and elections [\[51\]](#).

Several open questions remain to be explored. Given that smartphones provide easy access to a large range of apps, do certain types of exposure, such as positive versus negative content or news versus social media exposure, lead to more discussions in communication apps? Are individuals simply not encountering political content, or are they actively avoiding it? Answering these questions will provide insights into the evolving landscape of political engagement and the state of American democracy.

References

- [1] Michael X Delli Carpini and Scott Keeter. *What Americans know about politics and why it matters*. Yale University Press, 1996.
- [2] Jennifer Jerit, Jason Barabas, and Toby Bolsen. Citizens, knowledge, and the information environment. *American Journal of Political Science*, 50(2):266–282, 2006.
- [3] Robert A Dahl. *Democracy and its Critics*. Yale university press, 2008.
- [4] Nic Newman, A Ross Arguedas, Craig T Robertson, Rasmus Kleis Nielsen, and Richard Fletcher. *Digital news report 2025*. Reuters Institute for the study of Journalism, 2025.
- [5] Maria Petrova, Ananya Sen, and Pinar Yildirim. Social media and political contributions: The impact of new technology on political competition. *Management Science*, 67(5):2997–3021, 2021.
- [6] Erika Franklin Fowler, Michael M Franz, Gregory J Martin, Zachary Peskowitz, and Travis N Ridout. Political advertising online and offline. *American Political Science Review*, 115(1): 130–149, 2021.
- [7] Ian Vandewalker, Petry Eric, Brendan Glavin, Erika Franklin Fowler, Michael Franz, Pavel Oleinikov, Breeze Floyd, Travis Ridout, Yujin Kim, and Meiqing Zhang. Online ad spending in 2024 election totaled at least \$1.9 billion. *Brennan Center for Justice, Open Secrets, and Wesleyan Media Project*, 2025. URL <https://www.brennancenter.org/our-work/analysis-opinion/online-ad-spending-2024-election-totaled-least-19-billion>.
- [8] Markus Prior. *Post-Broadcast Democracy: How Media Choice Increases Inequality in Political Involvement and Polarizes Elections*. Cambridge University Press, 2007.
- [9] Sinan Aral. *The Hype Machine: How Social Media Disrupts Our Elections, Our Economy, and Our Health—And How We Must Adapt*. Crown Currency, 2021.
- [10] Soroush Vosoughi, Deb Roy, and Sinan Aral. The spread of true and false news online. *Science*, 359(6380):1146–1151, 2018.
- [11] Luca Braghieri, Sarah Eichmeyer, Ro’ee Levy, Markus M. Mobius, Jacob Steinhardt, and Ruiqi Zhong. Article-level slant and polarization of news consumption on social media. 2025.
- [12] Eli Pariser. *The Filter Bubble: How the New Personalized Web Is Changing What We Read and How We Think*. Penguin, 2011.
- [13] Jack Gillum, Alexa Corse, and Adrienne Tong. X algorithm feeds users political content—whether they want it or not. *Wall Street Journal*, 2024. URL <https://www.wsj.com/politics/elections/x-twitter-political-content-election-2024-28f2dadd>.

- [14] R Treisman. Meta is limiting how much political content users see. here’s how to opt out of that. npr, 2024.
- [15] Markus Prior. Who watches presidential debates? measurement problems in campaign effects research. *Public Opinion Quarterly*, 76(2):350–363, 2012.
- [16] Markus Prior. The challenge of measuring media exposure: Reply to dilliplane, goldman, and mutz. *Political Communication*, 30(4):620–634, 2013.
- [17] Seth Flaxman, Sharad Goel, and Justin M. Rao. Filter bubbles, echo chambers, and online news consumption. *Public Opinion Quarterly*, 80(S1):298–320, 2016. doi: 10.1093/poq/nfw006.
- [18] Ryan C Moore, Ross Dahlke, and Jeffrey T Hancock. Exposure to untrustworthy websites in the 2020 us election. *Nature Human Behaviour*, 7(7):1096–1105, 2023.
- [19] Sandra González-Bailón, David Lazer, Pablo Barberá, Meiqing Zhang, Hunt Allcott, Taylor Brown, Adriana Crespo-Tenorio, Deen Freelon, Matthew Gentzkow, Andrew M. Guess, et al. Asymmetric ideological segregation in exposure to political news on facebook. *Science*, 381(6656):392–398, 2023.
- [20] Byron Reeves, Thomas Robinson, and Nilam Ram. Time for the human screenome project. *Nature*, 577(7790):314–317, 2020.
- [21] Byron Reeves, Nilam Ram, Thomas N. Robinson, James J. Cummings, C. Lee Giles, Jennifer Pan, Agnese Chiatti, M. J. Cho, Katie Roehrick, Xiao Yang, et al. Screenomics: A framework to capture and analyze personal life experiences and the ways that technology shapes them. *Human–Computer Interaction*, 36(2):150–201, 2021.
- [22] Daniel Muise, David Markowitz, Byron Reeves, Nilam Ram, and Thomas Robinson. (mis)measurement of political content exposure within the smartphone ecosystem: Investigating common assumptions. *Journal of Quantitative Description: Digital Media*, 4, 2024.
- [23] Jason Barabas, Jennifer Jerit, William Pollock, and Carlisle Rainey. The question (s) of political knowledge. *American Political Science Review*, 108(4):840–855, 2014.
- [24] Charles Angelucci and Andrea Prat. Is journalistic truth dead? measuring how informed voters are about political news. *American Economic Review*, 114(4):887–925, 2024.
- [25] Jennifer Allen, Baird Howland, Markus Mobius, David Rothschild, and Duncan J. Watts. Evaluating the fake news problem at the scale of the information ecosystem. *Science Advances*, 6(14):eaay3539, 2020.
- [26] Andrew M Guess. (almost) everything in moderation: New evidence on americans’ online media diets. *American Journal of Political Science*, 65(4):1007–1022, 2021.
- [27] Ro’ee Levy. Social media, news consumption, and polarization: Evidence from a field experiment. *American Economic Review*, 111(3):831–870, 2021.

- [28] Eytan Bakshy, Solomon Messing, and Lada A. Adamic. Exposure to ideologically diverse news and opinion on facebook. *Science*, 348(6239):1130–1132, 2015.
- [29] Andrew M. Guess, Neil Malhotra, Jennifer Pan, Pablo Barberá, Hunt Allcott, Taylor Brown, Adriana Crespo-Tenorio, Drew Dimmery, Deen Freelon, Matthew Gentzkow, et al. How do social media feed algorithms affect attitudes and behavior in an election campaign? *Science*, 381(6656):398–404, 2023.
- [30] Solomon Messing. Are algorithmic bias claims supported? *Science*, 381(6665):1420–1420, 2023.
- [31] Nicholas Freiling. Just 37% of americans can name their representative. *Haven Insights*, 2017.
- [32] Shiro Kuriwaki. Cumulative ces common content, 2025. Cumulative Common Content dataset combining 2006–2024 Cooperative Congressional Election Study (CCES) / Cooperative Election Study (CES).
- [33] Caroline Le Pennec and Vincent Pons. How do campaigns shape vote choice? multicountry evidence from 62 elections and 56 tv debates. *The Quarterly Journal of Economics*, 138(2): 703–767, 2023.
- [34] Emily J Bell, Taylor Owen, Peter D Brown, Codi Hauka, and Nushin Rashidian. The platform press: How silicon valley reengineered journalism. 2017.
- [35] Cass R Sunstein. Republic: Divided democracy in the age of social media. 2017.
- [36] George Beknazar-Yuzbashev, Rafael Jiménez Durán, Jesse McCrosky, and Mateusz Stalinski. Toxic content and user engagement on social media: Evidence from a field experiment. 2022.
- [37] Hunt Allcott and Matthew Gentzkow. Social media and fake news in the 2016 election. *Journal of Economic Perspectives*, 31(2):211–236, 2017.
- [38] Stefano DellaVigna, Devin Pope, and Eva Vivalt. Predict science to improve science. *Science*, 366(6464):428–429, 2019.
- [39] Nir Grinberg, Kenneth Joseph, Lisa Friedland, Briony Swire-Thompson, and David Lazer. Fake news on twitter during the 2016 us presidential election. *Science*, 363(6425):374–378, 2019.
- [40] World Bank. Gini index for the united states. Retrieved from FRED, Federal Reserve Bank of St. Louis, 2025. <https://fred.stlouisfed.org/series/SIPOVGINIUSA>.
- [41] Lucas Chancel, Thomas Piketty, Emmanuel Saez, and Gabriel Zucman. *World inequality report 2022*. Harvard University Press, 2022.
- [42] Bo Cowgill. Economics, fairness and algorithmic bias. 2019.
- [43] Amanda Y Agan, Diag Davenport, Jens Ludwig, and Sendhil Mullainathan. Automating automaticity: How the context of human choice affects the extent of algorithmic bias. Technical report, National Bureau of Economic Research, 2023.

- [44] Sinan Aral and Paramveer S Dhillon. Digital paywall design: Implications for content demand and subscriptions. *Management Science*, 67(4):2381–2402, 2021.
- [45] Felix Chopra, Ingar Haaland, Fabian Roeben, Christopher Roth, and Vanessa Sticher. News customization with ai. 2025.
- [46] Romain Ferrali, Guy Grossman, and Horacio Larreguy. Can low-cost, scalable, online interventions increase youth informed political participation in electoral authoritarian contexts? *Science Advances*, 9(26):eadf1222, 2023.
- [47] Filipe R Campante, Ruben Durante, Felix Hagemeister, and Ananya Sen. Genai misinformation, trust, and news consumption: Evidence from a field experiment. Technical report, National Bureau of Economic Research, 2025.
- [48] John M. Abowd, Francis Kramarz, and David N. Margolis. High wage workers and high wage firms. *Econometrica*, 67(2):251–333, 1999.
- [49] Patrick Kline, Raffaele Saggio, and Mikkel Sølvsten. Leave-out estimation of variance components. *Econometrica*, 88(5):1859–1898, 2020.
- [50] Kari Paul. Elon musk reportedly forced twitter algorithm to boost his tweets after super bowl flop. *The Guardian*. <https://www.theguardian.com/technology/2023/feb/15/elon-muskchanges-twitter-algorithm-super-bowl-slump-report>, 2023.
- [51] Sinan Aral and Dean Eckles. Protecting elections from social media manipulation. *Science*, 365(6456):858–861, 2019.
- [52] Elliott Ash and Stephen Hansen. Text algorithms in economics. *Annual Review of Economics*, 15(1):659–688, 2023.
- [53] Giulia Caprini. Visual bias. 2023.
- [54] David M. Blei, Andrew Y. Ng, and Michael I. Jordan. Latent dirichlet allocation. *Journal of Machine Learning Research*, 3(Jan):993–1022, 2003.
- [55] Matthew Gentzkow, Bryan Kelly, and Matt Taddy. Text as data. *Journal of Economic Literature*, 57(3):535–574, 2019.
- [56] Reinhold Kesler, Michael Kummer, and Patrick Schulte. Competition and privacy in online markets: Evidence from the mobile app industry. In *Academy of Management Proceedings*, volume 2020, page 20978. Academy of Management, 2020.
- [57] Reinhold Kesler. The impact of apple’s app tracking transparency on app monetization. 2023.
- [58] Pauline Affeldt and Reinhold Kesler. Competitors’ reactions to big tech acquisitions: Evidence from mobile apps. 2021.
- [59] Rebecca Janssen, Reinhold Kesler, Michael E. Kummer, and Joel Waldfogel. Gdpr and the lost generation of innovative apps. Technical report, National Bureau of Economic Research, 2022.

- [60] Simon Kemp. Digital around the world. Available at <https://datareportal.com/global-digital-overview>, 2024.
- [61] Guy Aridor. Measuring substitution patterns in the attention economy: An experimental approach. *The RAND Journal of Economics*, 56(3):302–324, 2025.
- [62] Levi Boxell and Jacob Conway. Journalist ideology and the production of news: Evidence from movers. 2022.
- [63] Julia Cagé, Moritz Hengel, Nicolas Hervé, and Camille Urvoy. Political bias in the media—evidence from the universe of french broadcasts, 2002–2020. Technical report, CESifo Working Paper, 2025.
- [64] David Card, Ana Rute Cardoso, Joerg Heining, and Patrick Kline. Firms and labor market inequality: Evidence and some theory. *Journal of Labor Economics*, 36(S1):S13–S70, 2018.
- [65] Stéphane Bonhomme, Thibaut Lamadon, and Elena Manresa. A distributional framework for matched employer employee data. *Econometrica*, 87(3):699–739, 2019.
- [66] John M. Abowd, Robert H. Creecy, and Francis Kramarz. Computing person and firm effects using linked longitudinal employer-employee data. Technical report, Center for Economic Studies, US Census Bureau, 2002.
- [67] John M. Abowd, Francis Kramarz, Paul Lengermann, and Sébastien Pérez-Duarte. Are good workers employed by good firms? a test of a simple assortative matching model for france and the united states. 2004.
- [68] Rachel Hartman, Aaron J Moss, Shalom Noach Jaffe, Cheskie Rosenzweig, Leib Litman, and Jonathan Robinson. Introducing connect by cloudresearch: Advancing online participant recruitment in the digital age. 2023.
- [69] Benjamin Enke, Thomas Graeber, and Ryan Oprea. Confidence, self-selection, and bias in the aggregate. *American Economic Review*, 113(7):1933–1966, July 2023. doi: 10.1257/aer.20220915.
- [70] Manasi Deshpande and Rebecca Dizon-Ross. The (lack of) anticipatory effects of the social safety net on human capital investment. *American Economic Review*, 113(12):3129–3172, December 2023. doi: 10.1257/aer.20230010.
- [71] Stefano DellaVigna, Woojin Kim, and Elizabeth Linos. Bottlenecks for evidence adoption. *Journal of Political Economy*, 132(8):2748–2789, August 2024. doi: 10.1086/729447.

Acknowledgements

We thank conference/seminar audiences at Bocconi, CEPR Text-as-Data Workshop, Chicago Stigler Center, Econometric Society AMES and Europe Winter Meetings, HBS Digital Tech Regulation and Competition Conference, McGill University, NBER Summer Institute (Digital Economics and AI), Northwestern Kellogg, NYC Media Seminar, Ofcom, Paris Dauphine, Polarization in the Age of AI and the Post-Truth Era Workshop (Johns Hopkins), Quantitative Marketing and Economics Conference, Tufts University, University of Bern, UIUC, USC, and the Venice Summer Institute for helpful comments. We thank Annalí Casanueva Artís, Luca Braghieri, Jacob Conway, Matthew Gentzkow, Simona Mandile, Yotam Margalit, Melanie Monastirsky, Brendan Nyhan, Chris Roth, Muly San, and Andrey Simonov for helpful feedback and suggestions. We thank Daniel Muise and Justin Cornelius for assistance with Screenlake data, Brendan Nyhan for sharing the expert survey with the POLMETH mailing list, Reinhold Kesler for generously sharing Google Play Store data, and Bharad Raghavan for assistance in extracting time usage data. We thank Yunshu Cao, Nicola de Martini, Rotem Naimi, Michael Reeve, and Yuqin Wan for excellent research assistance. All errors are our own. The research in this article was approved by the Institutional Review Boards at Tel Aviv University (0009288-1) and Northwestern University (STU00224667).

Funding

Funding for this project was provided by the U.S.-Israel Binational Science Foundation (BSF) and the Tel Aviv University Center for AI & Data Science (TAD).

Author contributions

All authors contributed to the design, analysis, and writing of the research.

Competing interests

The authors declare no competing interests.

Omitted Figures and Tables

Figure 6: **Election-related exposures peak after polls close on election day.** The figure plots election-related exposure over time, showing the mean number of election-related seconds of exposure per hour across all active devices during the 48-hour window surrounding Election Day (November 5, 2024). The figure shows that election-related exposures were elevated throughout the day of the election, but that interest peaked around 8 PM when polls closed or shortly afterward (the vast majority of polls close at 7 or 8 PM local time). All timestamps are shown in Eastern Standard Time (EST). Vertical dashed lines indicate key election-night milestones. We use the time zone data provided by Screenlake, which is inferred from the various brands, apps, and/or terms that appear on screen.

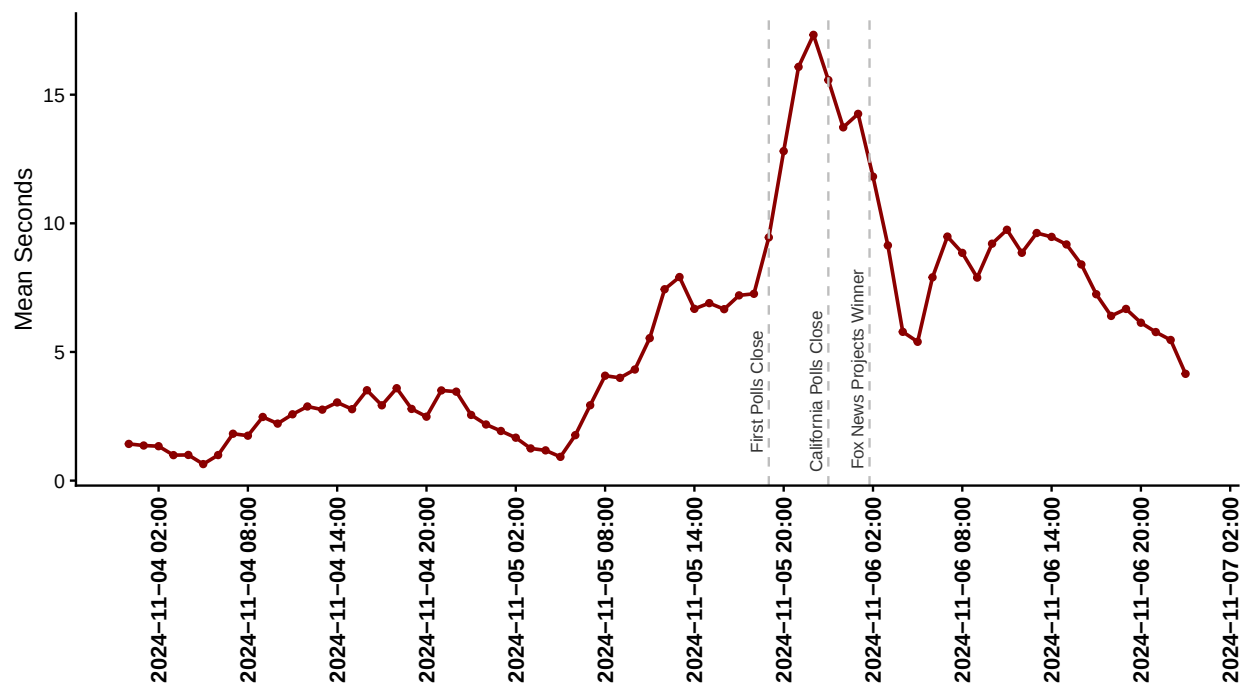
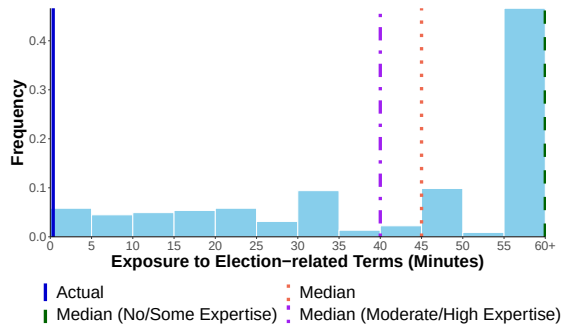
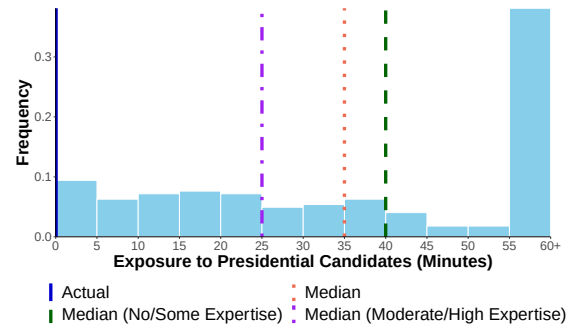


Figure 7: **Expert predictions diverge sharply from actual consumption levels.** Each panel presents the share of respondents predicting a given range of values (bars), alongside the median prediction for all respondents (coral dotted line), the median prediction for respondents with no or some expertise (green dashed line), the median prediction for respondents with moderate or high expertise (purple dot-dashed line), and the actual value (blue solid line).

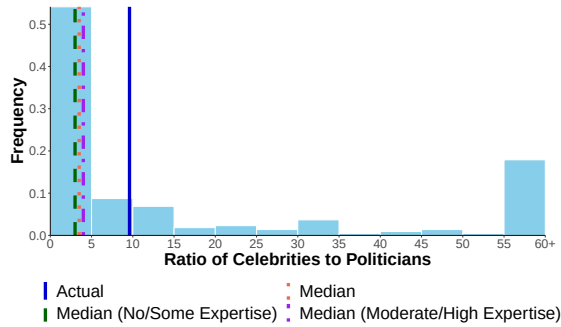
(a) Exposure to Election-related Terms (Minutes)



(b) Exposure to Presidential Candidates (Minutes)



(c) Ratio of Celebrities to Politicians



(d) Percentage of News Exposure (%)

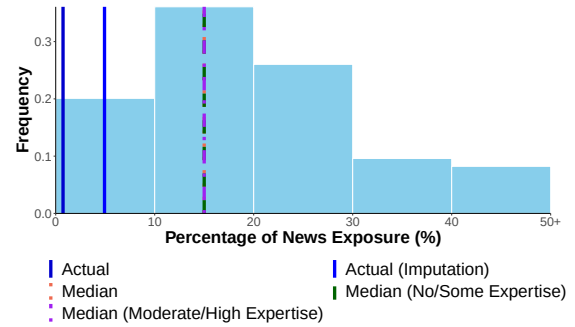


Figure 8: **Inequality in election-related exposure exceeds income inequality.** This figure plots Lorenz curves for time spent on the phone, time spent consuming TV, and smartphone exposures to election-related content. To calculate time spent on the phone, we use Screenlake application usage data from December 19th, 2024 until January 25th, 2025 since it is reliable for this time period. To calculate time spent consuming TV we use 2024 data from the American Time Use Survey (variable 120303).

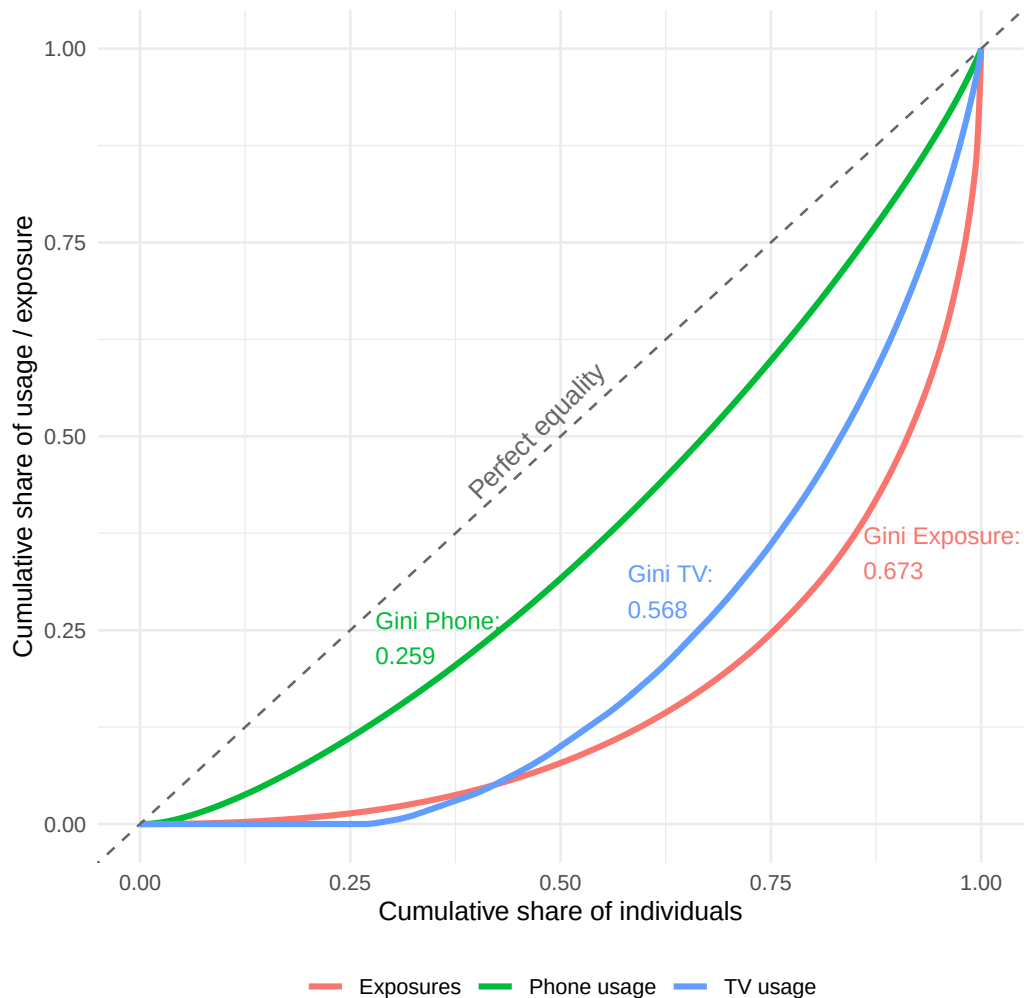


Figure 9: Election-related exposure varies across subgroups and applications. This figure demonstrates the variation in exposure to election-related content across individuals and applications. Panel (a) shows the distribution of average daily exposure (in seconds) of the median individual to election-related content across subgroups (September 1–November 4, 2024). Dots represent group medians; bars indicate interquartile ranges (25th–75th percentile). Subgroup labels include the group’s share of the total sample in parentheses. Geographic groups were assigned based on the overrepresented state method in encounter data (see Supplementary Section S7.1). Panel (b) shows variation in the share of time users spent on election-related content across social media applications. The share is computed as estimated time spent observing election-related content based on exposures divided by estimated time usage on the application. The time usage is based on exposures to the full set of keywords and the imputation in Supplementary Section S7.2. Both panels reflect aggregate user behavior between September 1 and November 4, 2024.

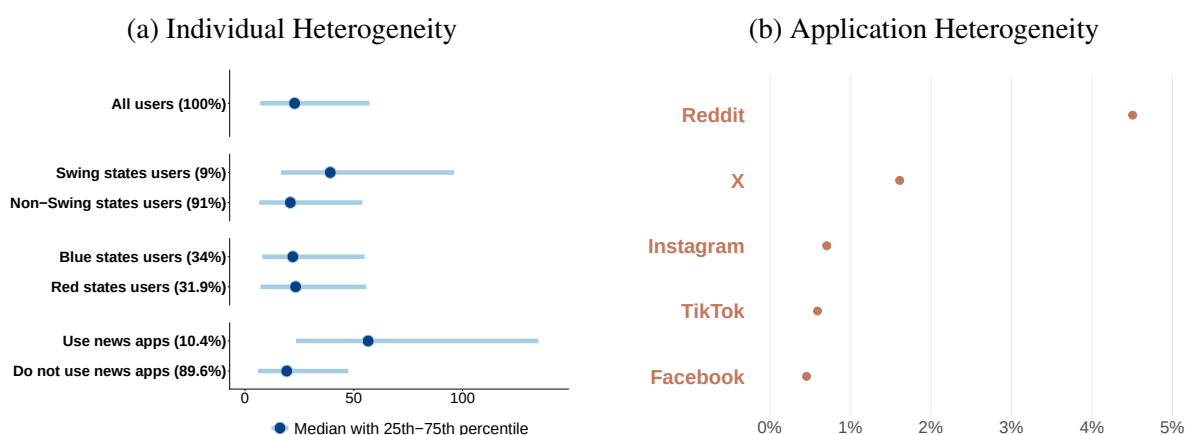


Figure 10: **X over-exposes users to election-related content and Elon Musk.** This figure presents the estimated coefficients of the app fixed effects from Equation (S1). We estimate this equation separately for each keyword group (election-related keywords, Elon Musk, and Mark Zuckerberg), restricting the analysis to the social media apps shown in the figure. We use data from September 1st, 2024 until November 4th, 2024. Facebook is the omitted category, serving as the reference group. The (small) 95% confidence intervals are constructed using standard errors clustered at the individual level. The dashed line in the top panel denotes the mean of the Instagram, Reddit, TikTok, and YouTube coefficients.

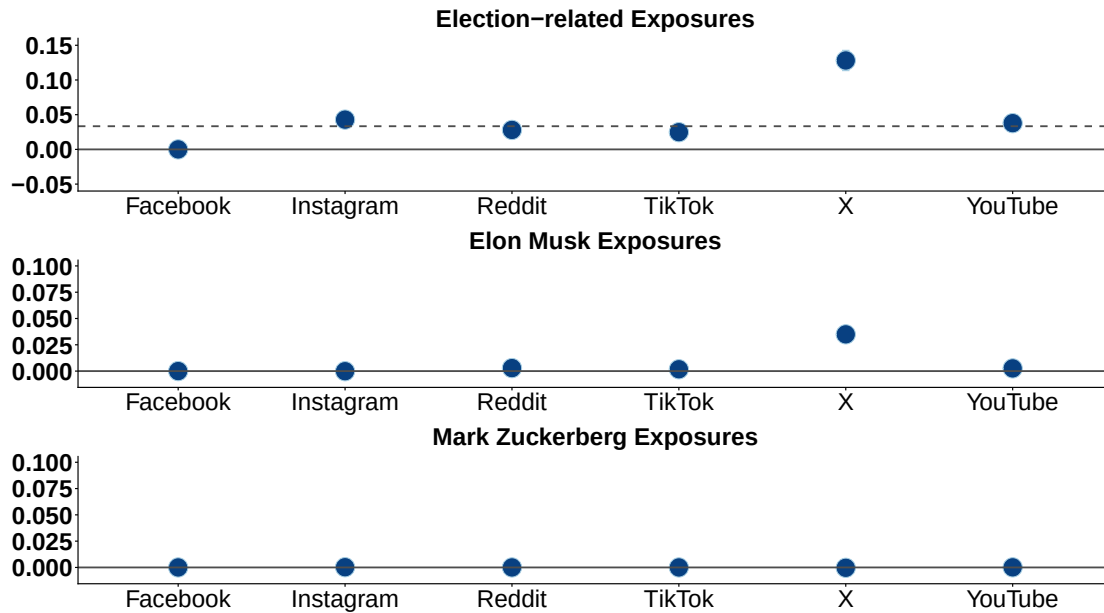


Table 1: **Median election-related consumption is robust across specifications.** This table reports the median number of election-related encounters and exposure in seconds across a range of robustness exercises. “Main analysis” corresponds to the baseline estimates reported in the main text. “With issues” expands the keyword set to include political issue terms (e.g., healthcare, immigration, inflation), counted only when appearing near other election-related words. “Reweight based on demographics” and “Reweight based on states” use weights to align the sample with the 2023 U.S. census distributions by age and gender, and the 2024 U.S. census distributions of red, blue, swing, and other states, respectively. “Exclude Spanish speakers” excludes users whose device language is set to Spanish. “Exclude non-election context” adjusts consumption using external precision and recall benchmarks from the New York Times benchmark described in Supplementary Section S1. “Assume two-second exposure” relaxes the main assumption that each exposure lasts three seconds. “Exclude first week” removes each user’s first week of data to address potential behavioral adaptation after the software installation.

	Encounters	Seconds
Main analysis	13.2	21
With issues	13.4	21.4
Reweight based on demographics	17.1	27
Reweight based on states	14.0	22.3
Exclude Spanish speakers	18.4	28.5
Exclude non-election context	7.8	12.4
Assume two-second exposure	13.2	14
Exclude first week	13.8	22

Table 2: **Variance decomposition of election-related keyword exposures (weekly).** This table presents the share of explained variance corresponding to each component of Equation (S2). We estimate this equation separately for each app category. The “All apps” column aggregates across all app categories: AI Assistant, Browser, Communication, Music & Video, News, Search, and Social Media. This aggregation is done by weighting the share of exposures of each group by the time spent on that group, using the category-specific median time frequencies from Supplementary Table S8. We exclude apps with fewer than 30 unique users and individuals with fewer than 30 unique apps.

	All apps	Social	Communication
Individual FE	0.78	0.70	0.83
App FE	0.22	0.28	0.11
Time FE	0.03	0.02	0.02
Covariance of ind. & app FE	-0.11	0.02	0.05
Covariance of ind. & time FE	0.00	-0.02	-0.02
Covariance of app & time FE	0.08	0.00	-0.00
Adj. R^2	0.26	0.27	0.21
N apps	285	10	13
N individuals	1469	1469	1469

Supplementary Information

S1 Keyword Informativeness Benchmarks

In this section, we provide a formulation of our measurement problem as a standard topic classification problem. Our goal is to analyze whether our keywords are predictive of election-related content. Let t denote a time period within a given day and assume that an individual encounters content $\mathbf{w}_t = (w_{t,1}, \dots, w_{t,N_t})$ in period t (dropping individual subscripts for simplicity). Following standard practice in the natural language processing (NLP) literature [52], we represent content as a sequence of tokens, which can represent words but also more general objects such as punctuation or emojis.¹⁰ Let $\mathcal{V}$ denote the universe of tokens that individuals encounter on their phone, and $\mathcal{K} \subset \mathcal{V}$ denote the set of election-related keywords that are tracked. We use the number $k \in \{1, \dots, |\mathcal{K}|\}$ to refer to keywords in that set.

We define a set of topics $\mathcal{T}$ and denote the set of topics that content $\mathbf{w}_t$ relates to by $topics(\mathbf{w}_t) \subset \mathcal{T}$.¹¹ Thus, we say that content is election-related when the US 2024 Election is in the set of topics of that content, or $election \in topics(\mathbf{w}_t)$. To ease notation, we define the “ground-truth” indicator of whether the content encountered by an individual on t is election-related:

$$E_t^* = \begin{cases} 1, & \text{if the individual consumed election content in period } t, \\ 0, & \text{otherwise.} \end{cases}$$

We also define our empirical building block: an exposure, defined as an indicator of whether keyword k is displayed in period t , $d_t(k) \in \{0, 1\}$. Our keyword-based proxy of E_t^* , is thus:

$$E_t(\mathcal{K}) = \begin{cases} 1, & \text{if at least one keyword } k \text{ is displayed in period } t, \\ 0, & \text{otherwise.} \end{cases}$$

Because ultimately our goal is to measure E_t^* (true consumption) using $E_t(\mathcal{K})$ (observed keywords), we focus on two performance metrics that are standard in text-based classification: *precision* and *recall*. Precision is the probability that content is election-related given exposure to one of our election-related keywords, $\Pr(E_t^* = 1 | E_t(\mathcal{K}) = 1)$, while recall is the probability that one of our keywords is detected when the content is election-related: $\Pr(E_t(\mathcal{K}) = 1 | E_t^* = 1)$. We

¹⁰It need not be limited to textual content and can also be textual representations of images or videos [53].

¹¹Topics are canonically defined in the NLP literature as a probability vector over possible tokens [54], allowing content to be related to multiple topics. However, the literature also often assumes that content has a single topic [55]. We can relate our *topics* function to the literature by assuming that it extracts the most likely topics of a given sequence of tokens (e.g., with probability above a certain threshold) or their modal topic.

emphasize that this represents the exposure measure from Section 1.

A challenge in our setting is that our keywords had to be pre-specified, preventing us from choosing keywords ex-post based on F1 scores that trade off between precision and recall. Nevertheless, we pre-specified a set of keywords that we ex-ante believed were predictive of election-related content. Below, we confirm with several metrics using different definitions of ground truth that this is indeed the case. In particular, we collect the universe of articles published by a left-leaning (the New York Times) and a right-leaning (Fox News) news organization during September and October 2024. We then emulate the Screenlake keyword detection algorithm and measure keyword occurrences in the way that they would be logged by the software.

Defining Ground-Truth Topic Labels. We rely on the labels that the New York Times and Fox News place on their articles in order to determine whether an article is election-related. The assumption is that keywords that are highly likely to be used to discuss the election are present in articles tagged as ‘elections’ or ‘politics’ on the New York Times or ‘politics’ or ‘official-polls’ on Fox News. To define articles that are *not* election-related, we consider article tags that are probably not related to politics but would be discussed in the news, such as cultural topics, sports, and weather. We then consider the articles within the following set of article tags as our ‘non-election’ articles: ‘Arts’, ‘Style’, ‘Books’, ‘Movies’, ‘Food’, ‘Real Estate’, ‘Climate’, ‘Well’, ‘Technology’, ‘Science’, ‘Health’, ‘Weather’, ‘Theater’, ‘Travel’, ‘Sports’ on the New York Times; and ‘Sports’, ‘Travel’, ‘Food-Drink’, ‘Lifestyle’, ‘Entertainment’, ‘Tech’, ‘Health’, ‘True-Crime’, ‘Science’, ‘Family’, ‘Faith-Values’, and ‘Weather’ on Fox News.¹² We drop articles that do not have any of these tags (as it is ambiguous if they are election-related) or are shorter than 100 words.

Precision and Recall in News Benchmarks. We find that recall is high at 0.98 for both benchmarks. Specifically, 97.74% and 97.89% of articles about the elections in the New York Times and Fox News, respectively, contained at least one of our keywords. If we consider 250-word increments of these articles, a rough approximation to the content visible on the screen at a given time, our recall is 0.87 and 0.88 on the New York Times and Fox News, respectively. This means that if someone saw even part of the article on their screen, one of our keywords would likely be detected.

We then compute the precision statistic for both benchmarks and find that it is 0.58 and 0.70 for the New York Times and Fox News, respectively. Additionally, we find that our keywords cover 7 of the top 10 words with the highest F1 score across both Fox News and New York Times, indicating that, even though we necessarily had to pick our keyword set before the election, it is

¹²Note that any election-related articles that may appear within these topics, for example in climate, will artificially deflate our precision score.

still highly informative ex-post and among the most informative set we could have selected.¹³

Recall in User-Generated Content. One concern with using news articles as a benchmark for computing precision and recall is that the language that individuals use to discuss elections may be different than the language written by journalists and edited by news editors. In order to construct a proxy for recall in the context of user-generated content, we pull the full set of posts from the ‘r/politics’ subreddit during September and October 2024. Mirroring the full article as the unit of analysis in the news benchmarks, we use the full set of comments under each post and compute our recall statistic over this sample. We find that 90% of the posts have at least one of our keywords in their comment thread, which remains similar once we weight posts by a proxy of their popularity (upvotes on Reddit).

To summarize, we make two main points. First, we provide evidence that our keyword-based metric of exposure achieves a high recall and moderate precision, using labels from New York Times and Fox News articles as ground truth. This means that, while not exhaustive, exposures to our primary keyword set are highly likely to provide an *upper bound* of overall consumption of election-related content from traditional media. Second, we confirm that this set also has high recall of election-related content that is user-generated.

S2 External Validity

In this section, we assess external validity. This is important as our sample is composed of individuals who opt to use the application offered by Screenlake. There are three key concerns for how the results from our sample may not generalize to the broader U.S. population that we address.

1. **Demographic representativeness:** The Screenlake sample may not be demographically representative.

How we address it: We apply Screenlake-provided demographic weights to match a representative age–gender distribution based on the 2023 U.S. census. We replicate our main figures and results using these weights and find that the results do not qualitatively change. We also conduct a similar exercise, reweighting our sample to match the distribution of red, blue, and swing states.

2. **Different phone usage patterns:** Screenlake users may differ from the typical American in the apps they install and how they use their phones.

¹³The top 10 words with the highest F1 score are: ‘trump’, ‘president’, ‘harris’, ‘election’, ‘campaign’, ‘former’, ‘kamala’, ‘vice’, ‘biden’, ‘donald’.

How we address it: We benchmark (i) installed applications against market-level Google Play download data and (ii) time on device and by app category against both industry benchmarks and a U.S.-representative Prolific sample. We find that our sample uses a similar set of applications and has similar overall engagement levels to the representative benchmarks.

3. **Screen-time management motivation:** Individuals installing screen management apps may have different news consumption propensity.

How we address it: We conduct a survey using Prolific-recruited participants asking whether individuals use screen-time applications and their self-reported propensity to consume news. We find little difference in self-propensity to consume news across this dimension.

We now document the details for each of these exercises that lead us to conclude that our main results are likely robust and generalizable to the broader U.S. population.

S2.1 User Demographics

First, we consider the demographic composition of our sample using individual-level weights provided by Screenlake. The data collected by Screenlake is anonymous and disconnected from identifiers. Screenlake infers gender and age post-collection based on the inter-demographic relative likelihood of combinations of app usage and various terms appearing on-screen. Our sample skews male and younger than the general population (over 60% men and over half under the age of 35), and therefore, we analyze how our results change when reweighting the sample. Screenlake provided us with individual-level weights, making our sample representative of the U.S. population on age and gender, according to 2023 U.S. Census data.

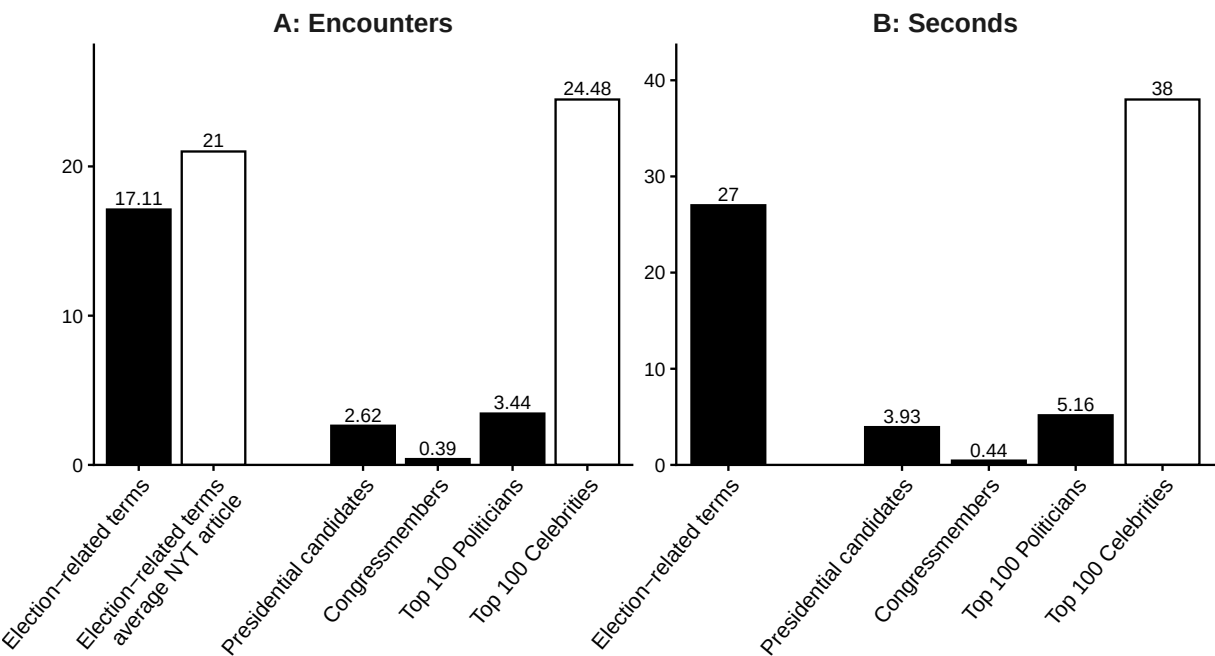
We replicate our main results using the reweighted sample. Supplementary Figure S1 replicates Figure 1 and shows that the median exposures and encounters of average daily election-related content consumption increase but remain qualitatively low. For example, we find that the exposure ratio between the top 100 celebrities and politicians remains high at 7.4. Supplementary Figure S2 analyzes exposure over time with the reweighted sample and presents a similar pattern to Figure 2. Finally, we compute the percentage of total election-related content consumption coming from news applications (replicating the main result of the “Sources of Information” subsection) and find that it still remains low—0.89% and 0% for the mean and median individual, respectively. Overall, we conclude that our key results on the magnitude, timing, and source of election-related content consumption are qualitatively robust to demographic reweighting.

Additionally, we reweight our sample based on the location of individuals to account for potential nonrepresentativeness in location. Specifically, the reweighting is performed at the level of

four political state groups: blue states, red states, swing states, and other states, where we cannot distinguish between two states (e.g., North and South Carolina). We reweight the sample to match the distribution of US population among these state groups based on 2024 U.S. census data. Supplementary Section S7.1 discusses the method we use to predict the state for each individual and provides the full list of states in each group. As shown in Extended Data Table 1, reweighting based on states results in very similar election-related consumption.

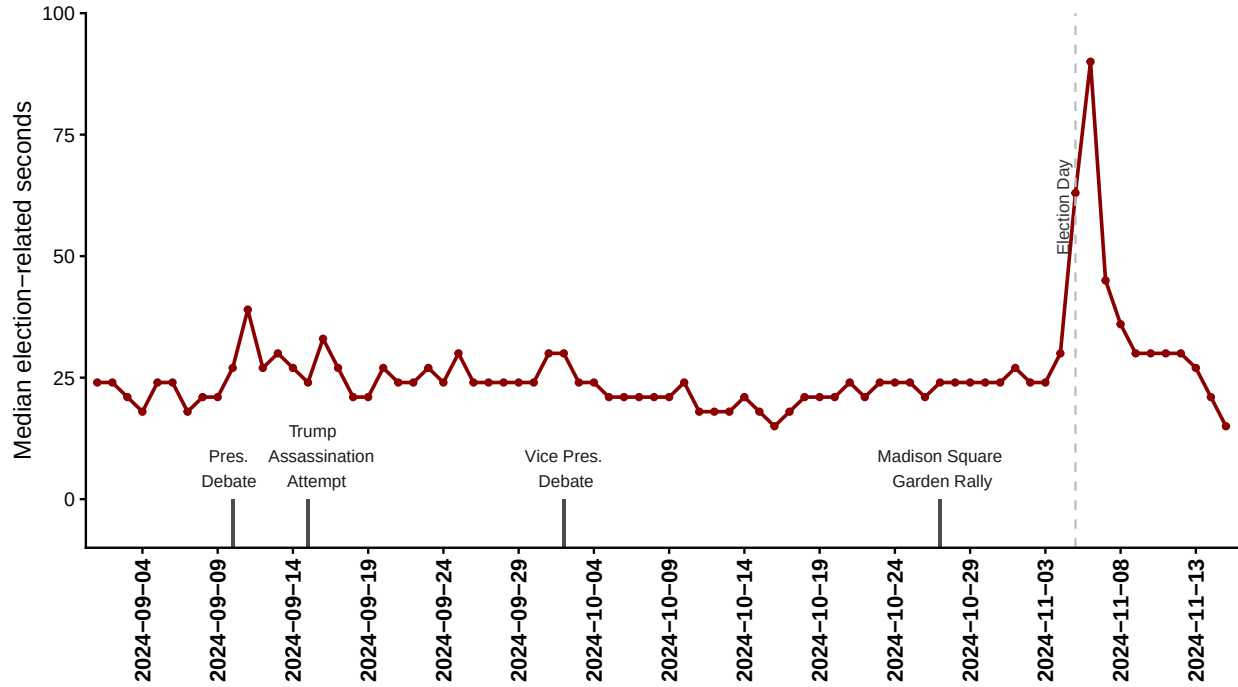
There is still a concern that the sample may differ from the population in unobserved ways. A particularly important concern for our main result is whether people in our sample systematically encounter less political content than the average person in the U.S. To understand selection into the sample, we complement our analysis with additional datasets.

Figure S1: Election-Related Content Consumption, Reweighted



NOTES: The figure reports the median encounters (panel A) and seconds of exposure (panel B) among active devices on an average day between September 1 and November 4, 2024. From left to right, each panel presents the overall consumption based on the full set of election-related terms, the number of election-related terms in an average New York Times article (only for encounters), names of presidential candidates, congressmembers, top 100 politicians, and top 100 celebrities. This figure uses demographic weights provided by our data provider that reweights our sample to be demographically representative on gender and age based on the 2023 U.S. census.

Figure S2: Political Exposures Over Time, Reweighted



NOTES: The figure displays the median daily number of seconds of exposure to election-related terms per active user between September 1 and November 15, 2024. Key campaign events are indicated along the x-axis: the Presidential Debate between Kamala Harris and Donald Trump (September 11), the second Trump assassination attempt (September 15), the Vice Presidential Debate between Tim Walz and J. D. Vance (October 2), and Trump’s Madison Square Garden rally (October 27). The vertical dashed line marks Election Day (November 5). This figure uses demographic weights provided by our data provider that reweights our sample to be demographically representative on gender and age based on the 2023 U.S. census.

S2.2 Phone Usage Patterns

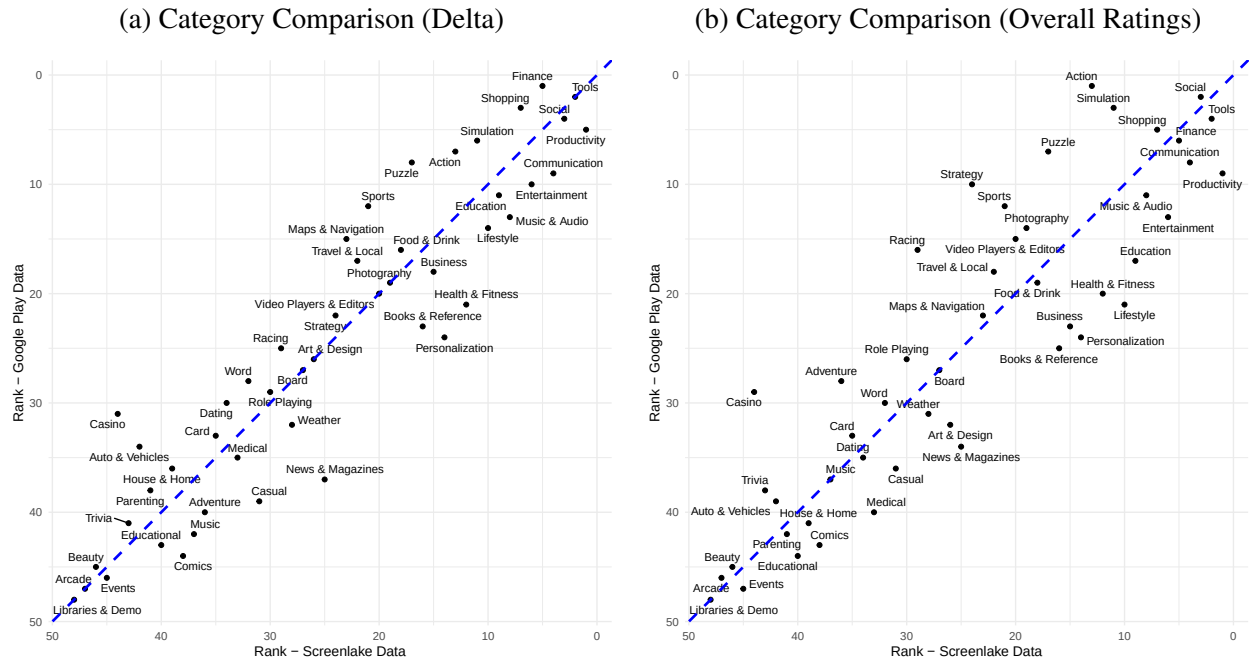
Comparison of Installations to Google Play Store Data. We now turn to assessing whether the Screenlake users have systematic differences in the types of applications that they have installed on their phones compared to a typical American. To do this, we compared the distribution of apps installed to population-level data available via the Google Play Store – the primary distribution channel for Android applications. We use data on the number of US ratings scraped during October 2024 and January 2025 for all applications with over 1 million downloads, or in the News & Magazines and Social categories.¹⁴ We follow the best practices for using this data and take the number of ratings for an application as an approximation for its overall popularity [56, 57, 58, 59].¹⁵ We consider two variants of this: total number of ratings (denoted overall rat-

¹⁴We thank Reinhold Kesler for generously sharing this data with us.

¹⁵Ratings are typically used since the publicly available range for the number of installations is too wide to provide a meaningful ordering of applications.

ings) measuring cumulative popularity and the difference in the number of ratings between the October 2024 and January 2025 data collection (denoted delta) measuring recent popularity.¹⁶ We compare the overall number of downloads to the overall prevalence of the same applications within our sample. As we do not directly observe installed applications for the individuals in our sample and only observe measures of time usage, we mark an application as being installed for an individual if we ever observe them using the application. Finally, in order to ensure an apples-to-apples comparison, we remove any applications that are pre-installed on Android phones and consider only the set of applications that we observe both in our dataset as well as in the Google Play data.¹⁷

Figure S3: Comparison between Screenlake Installed Applications and Google Play



NOTES: This figure compares the set of applications installed by the Screenlake panelists compared to market-level data from the Google Play Store. For the Google Play Store data, we proxy for total downloads using cumulative number of ratings in (b), and using the difference between the cumulative number of ratings from January 2025 and October 2024 in (a). For the Screenlake data, we mark a user as having an application installed if we ever observe them using it. Panels (a) and (b) provide a rank-rank plot of comparison across downloads aggregated to application categories.

Figures S3a and S3b show a remarkably strong correlation between apps used among our sample and the Google Play data, using the delta measure and overall ratings, respectively. Since there is wide heterogeneity in usage within categories across individuals, we compare the relative importance of categories (as defined by Google Play) across the two datasets. For both datasets, we sum

¹⁶For example, an application such as Among Us was very popular several years ago, and so accumulated a large number of ratings, but is not very popular during our sample period.

¹⁷Across all individuals in Screenlake data, we observe 26,015 unique applications.

across all of the individual applications within the dataset to obtain a total number of installations or ratings for the category and compare the relative rank across the datasets. These results show strong alignment in the relative set of installed applications across categories, especially according to the more reliable delta measure, indicating that the set of individuals in the sample is reasonably representative in terms of the overall category of applications they use. We compute the Spearman correlation coefficient, a commonly used measure of accordance for ranking data, and find that it is very strong, with 0.928 for the delta measure and 0.888 for the overall rating measures. Notably, according to this analysis, our sample has a *larger* relative importance of news applications compared to the general prevalence of these applications in the population.

We then turn to comparing the ranks of applications within our categories of interest for our primary analyses—news, communication, social, music & audio, and entertainment. The top 10 applications in terms of installations in our sample are: Instagram, Spotify, Facebook, TikTok, Telegram, Snapchat, Discord, Netflix, X, and Reddit. Similarly, the top 10 applications in terms of installations using the delta method in the Google Play sample are: Facebook, Instagram, Truecaller, TikTok, Snapchat, Facebook Lite, WhatsApp, Max, Telegram, and Spotify. This highlights substantial overlap in the top set of applications between our sample and the general population. Beyond the top applications, we consider the full set of applications that we observe across the categories of interest. We compute the weighted Spearman rank correlation coefficient, considering as the weight $\frac{1}{\text{Google Play Rank}}$ in order to place more weight on accordance for higher ranks, relative to those in the long tail. This provides a correlation coefficient of 0.896 for the delta measure and 0.920 for the overall ratings measure, indicating strong agreement in the relative rankings.

Overall, we conclude that the set of apps that the individuals in the sample use is reasonably representative of the app usage of the broader population.

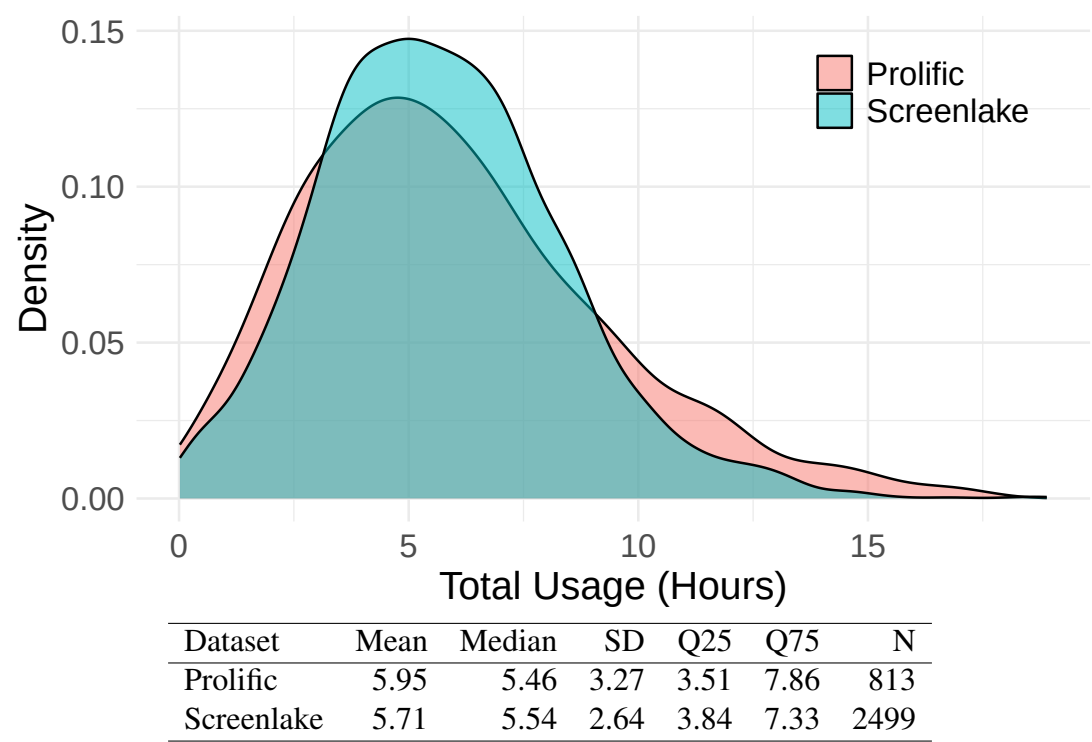
Comparison of Time Usage to Benchmarks and Representative Sample. While the set of installed apps among our sample could be reasonably representative, our sample may still differ from the general population in how often each app is used. Therefore, we compare the time usage of Screenlake participants to an industry benchmark and a representative sample of Prolific participants. We consider time usage observed in our dataset from December 19th, 2024 until January 25th, 2025 as this represents the time period where we have the most reliable measurement of application time usage.

We Are Social [60], which aggregates information from multiple sources, provides a useful industry benchmark. On average, Screenlake individuals spend slightly longer (3.2 hours a day) on a set of communication and social media apps than the average individual (2.5 hours a day).¹⁸

¹⁸The set of apps included in this benchmark are Facebook, Discord, Instagram, LinkedIn, Messenger, Pinterest, Reddit, Snapchat, Telegram, Threads, TikTok, WhatsApp, X, and YouTube.

Industry benchmarks may not correspond to the time period of analysis for our sample. To establish a representative baseline of smartphone usage for around the same period and study the time spent on specific apps, we recruited 813 Android individuals on Prolific from October 30 to November 3, 2024. The sample is representative in terms of average age and political affiliation. Participants are asked to complete a demographic survey and upload screenshots of their daily app-level screen time for the past 7 days using Digital Wellbeing, capturing “the entire list of apps that were used for more than 1 minute.”¹⁹

Figure S4: Distribution Comparison between the Representative Prolific and Screenlake



NOTES: This table presents the distribution of the average daily hours spent on the phone, per individual, for the representative Prolific sample and for Screenlake data, on which our sample is based. We use Screenlake application usage data from December 19th, 2024 until January 25th, 2025 since it is reliable for this time period.

Figure S4 shows that the distribution of overall phone usage is similar between our data (based on Screenlake) and the representative Prolific sample.²⁰ The median individual in both samples spends a strikingly similar amount of time on their phone, 5.46 and 5.54 hours, respectively, in

¹⁹In order to extract the time usage from the screenshots, we developed an OCR script using the OpenAI API.
²⁰We omit usage of Google Chrome from both samples, as Prolific users have a large amount of time on Google Chrome due to completing surveys for their work on Prolific in Chrome.

Table S1: Comparison of App Usage Between Prolific Data and Screenlake Data

	Average Hours Per Person		Average Hours If Used		Fraction of Nonzero Users		
	Prolific	Screenlake	Prolific	Screenlake	Prolific	Screenlake	Pew
Facebook	0.50	0.20	0.84	0.44	0.60	0.45	0.70
Instagram	0.21	0.63	0.44	0.96	0.47	0.66	0.50
LinkedIn	0.00	0.00	0.05	0.04	0.06	0.11	0.32
Reddit	0.22	0.04	0.50	0.24	0.44	0.17	0.24
Snapchat	0.05	0.07	0.23	0.24	0.20	0.28	0.27
TikTok	0.40	0.55	1.21	1.32	0.33	0.41	0.33
X	0.10	0.04	0.49	0.24	0.21	0.18	0.21
WhatsApp	0.06	0.41	0.38	0.58	0.16	0.70	0.30
Spotify	0.04	0.06	0.15	0.11	0.27	0.52	NA
YouTube	0.60	1.07	0.80	1.18	0.76	0.90	0.85
Fox News	0.00	0.00	0.13	0.00	0.00	0.00	NA
Google News	0.00	0.00	0.03	0.01	0.00	0.02	NA
NYTimes	0.00	0.00	0.25	0.06	0.01	0.01	NA
Netflix	0.04	0.06	0.51	0.27	0.08	0.20	NA

NOTES: This table presents the average daily hours spent on individual applications for the Prolific sample and for the Screenlake sample when an individual was active. We use Screenlake application usage data from December 19th, 2024 until January 25th, 2025 since it is reliable for this time period. The first two columns present the average number of hours spent on the given application per individual. The second two columns present the average number of hours spent on the given application per individual, conditional on using the application at all. The final three columns present the fraction of individuals who use the application in the Prolific sample, Screenlake sample, and from a PEW survey.

the Prolific representative sample and in our data (a 1% difference). The average individual in the Prolific sample spends 5.95 hours on their phone, while the average individual in our data spends 5.71 hours (a 4% difference). The results suggest that the Screenlake population (on which our data is based) is reasonably similar in overall phone usage to the representative sample from Prolific.

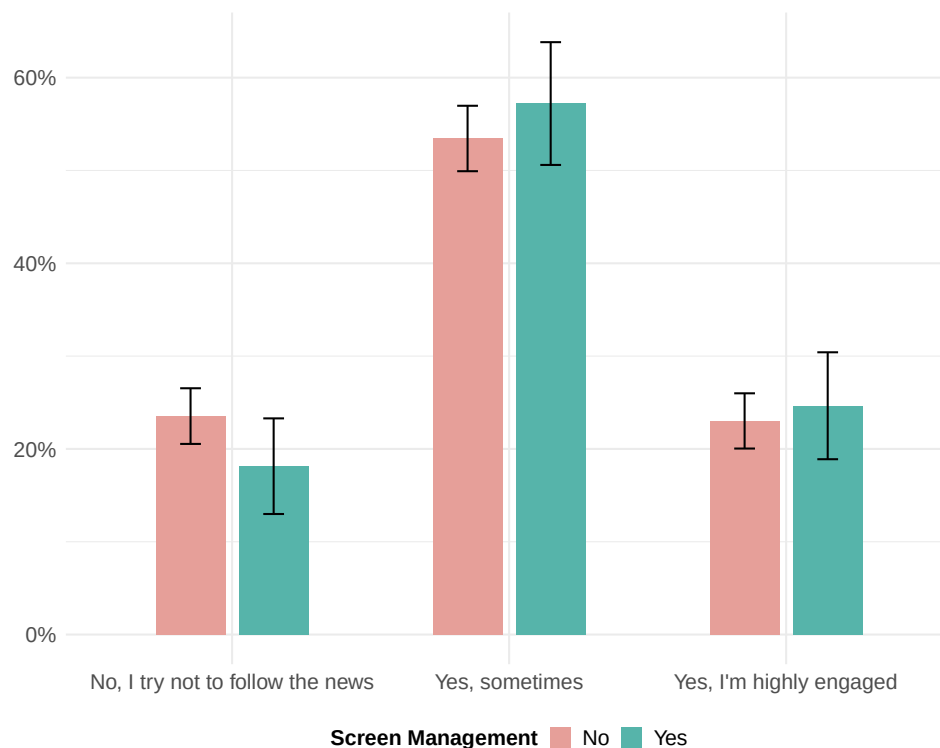
Next, we compare the time spent and usage of top applications – Facebook, Instagram, Netflix, Reddit, Snapchat, Spotify, TikTok, WhatsApp, X, and YouTube – across our data (based on Screenlake), the Prolific sample, and a sample from a Pew Research Center survey.²¹ The results are presented in Table S1. One notable difference between the samples is that the fraction of WhatsApp and Instagram usage for the Screenlake sample is higher than both the Prolific and Pew samples, while Facebook usage is lower. The Prolific sample has an over-representation of Reddit, both relative to Pew and Screenlake. Nonetheless, conditional on using the application, the two samples show similar patterns: large time spent on YouTube and TikTok, very similar time spent

²¹The Pew data is available at the extensive margin. See <https://www.pewresearch.org/internet/fact-sheet/social-media/> for details.

on Snapchat and Spotify, and reasonably similar time spent on X and WhatsApp. Their overall time spent on social media is around the same. Importantly, in both Prolific and Screenlake almost no time is spent on news apps, suggesting that the low share of exposure to election-related content from such apps is not due to differences between our sample and the general population. Overall, the Screenlake sample differs on several dimensions from these two benchmarks, but reflects similar engagement patterns across the core applications that individuals spend time on their phones.

S2.3 Comparison of Individuals with and without Screen Management Apps

Figure S5: Screen Management and Level of Political Engagement on Phone



NOTES: This figure presents the distribution of responses, partitioning on whether the respondents use screen management software and the frequency with which they follow the news on their phone. The denoted ranges represent 95% confidence intervals.

Finally, we assess whether there are systematic differences in self-reported news consumption between individuals who use screen management software and those who do not. As part of a news consumption survey,²² we asked respondents whether they use screen management software and

²²For more details of the survey, see Supplementary Section S5.4.

whether they try not to follow the news, they sometimes follow the news, or are highly politically engaged on their phones. We note that a relatively large fraction of the sample (21.8%) self-reports using screen management software, meaning that we have enough power to study the differences across the groups and that downloading these apps is relatively common.

Figure S5 presents the distribution of responses, partitioning on whether the respondents use screen management software, and finds little difference in political engagement between these sets of respondents. If anything, individuals with screen management software tend to be slightly more politically engaged. Hence, our results that individuals consumed arguably little news before the election are unlikely to be driven by the characteristics of people who download screen management apps.

S3 Additional Results for Election-Related Consumption

S3.1 Overall Consumption

In this section, we provide additional robustness exercises for our results measuring the median consumption levels of election-related content. Our qualitative conclusion of arguably low consumption does not vary across each of these exercises: reweighting to a representative sample, excluding Spanish speakers, considering an expanded keyword set that includes political issues, and making alternative measurement assumptions. We summarize the change in overall median consumption in Extended Data Table 1 and now discuss the details for each.

Reweighting. Our data provider provides a set of individual-level weights that allow us to reweight our sample to match the 2023 U.S. census on gender and age. Figure S1 shows that this consumption slightly increases but does not change dramatically after reweighting. In particular, reweighted median encounters and exposure seconds increase to 17.1 and 27, respectively. While the reweighted data shows slightly elevated levels, it does not change our qualitative conclusion that election-related consumption is small. Additionally, we reweight based on the location of individuals. We find that this reweighting keeps median encounters and exposure seconds almost the same as our main estimate, at 14.0 and 22.3, respectively. The reweighting exercises are described in more detail in Supplementary Section S2.1.

Excluding Spanish Speakers. Another potential concern is that our keyword set is in English, while in the U.S. there are many individuals whose primary language is Spanish. While we would still be able to detect names, we may not observe some of our election-related keywords for Spanish-speaking individuals. We can identify these individuals by looking for the presence of the application ‘Teléfono’ or ‘Llamar’ instead of ‘Phone’, meaning the individual probably has Spanish set as the default phone language. If we remove these individuals from our data, consump-

tion increases, but still remains low compared to the benchmarks discussed in the main text.

Expanded Keyword Set. We compute the overall median exposure to an expanded keyword set, with a set of political issues in addition to the election terms and politicians in our primary keyword set.²³ Since many of these terms are typically used in non-political contexts,²⁴ we only consider them as an election-related exposure if they are mentioned near a much larger set of words that could be related to the election. To provide additional context on the occurrence of a keyword, the SDK records the co-occurrence of a large list of surrounding words. In particular, upon observing the presence of a keyword on the screen, the application scans the ten observed words, excluding stopwords, before and after the identified keyword and logs all of the occurrences of surrounding words within this set. We are able to observe a large set of surrounding words related to the election, with the full list presented in Supplementary Section S9.1. This makes it so that these low precision keywords are converted into high precision encounters when paired with these surrounding words. So, for example, in the sentence: ‘the president announced a plan to improve healthcare’, we will consider the word ‘healthcare’ as a political issue keyword, since ‘healthcare’ is surrounded by ‘president’. In contrast, ‘healthcare’ will not be considered a political issue keyword in the sentence ‘are you pleased with your healthcare provider’. The median individual still only consumes approximately 21.4 seconds of election-related content after adding political issues.

Alternative Measurement Assumptions. We consider three additional computations of the median consumption that make alternative measurement assumptions.

First, we consider excluding non-election context. As noted in Supplementary Section S1, our keyword-based classifier has high recall but moderate precision. We can therefore compute a version of overall consumption where we divide by the recall in the New York Times benchmark and then multiply by the New York Times precision. This provides an alternative measure filtering out our false positives. Since we have high recall and moderate precision, this adjustment lowers our estimates for news consumption and results in 7.8 encounters for the median individual and 12.4 seconds of exposure.

Second, in our main analysis we conservatively assume that an exposure is logged as three seconds since we know that software will only capture keywords every three seconds. However, in reality this is an upper bound since most content is likely visible for a fraction of that time. Therefore, we consider a less conservative assumption of exposure seconds where we assume that an exposure corresponds to two, as opposed to three, seconds. This leads the median exposures to reduce to 14 seconds.

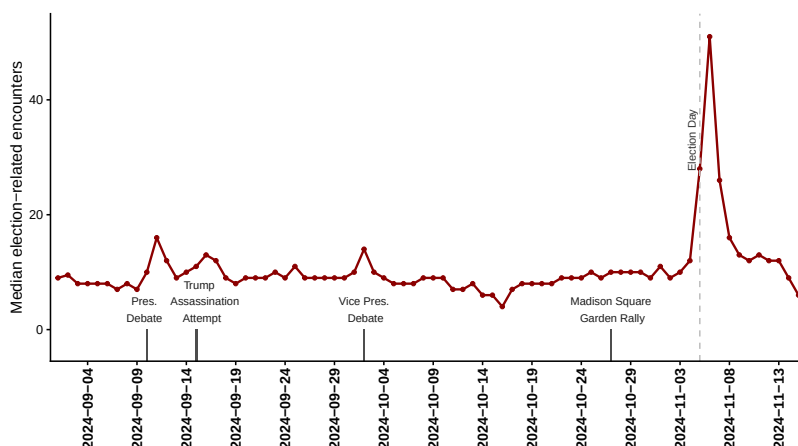
²³This list of keywords for the issues includes: “climate”, “court”, “crime”, “economy”, “healthcare”, “inflation”, “women”, “Gaza”, “West Bank”, “immigration”, “border”, “woke”, “DEI”, “Roe”, “gun”, “protest”, “recount”, “interference”, “hoax”.

²⁴We validated that these keywords have low precision using our New York Times articles discussed in Supplementary Section S1.

Third, as individuals may potentially change their behavior due to the screen management application being installed in the first week of usage, we recompute the same statistics but drop every individual's first week of usage. This barely affects our estimates and results in 13.8 encounters for the median individual and an exposure of 22 seconds.

S3.2 Consumption over Time

Figure S6: Median Daily Election-Related Encounters



NOTES: The figure displays the median daily number of election-related encounters per active user between September 1 and November 15, 2024. Key campaign events are indicated along the x-axis: the Presidential Debate between Kamala Harris and Donald Trump (September 11), the second Trump assassination attempt (September 15), the Vice Presidential Debate between Tim Walz and J. D. Vance (October 2), and Trump's Madison Square Garden rally (October 27). The vertical dashed line marks Election Day (November 5). See the Methods for the definition of election-related encounters.

S3.3 Sources of Information

Table S2: Share Exposed by Keyword Set and Category

Keyword Set	Category	Share Exposed	
		Average Day	Entire Period
Primary Political Terms	All	0.787	0.986
	News	0.024	0.090
	Social Media	0.403	0.763
	Browser	0.376	0.838
	Music and Video	0.316	0.807
	Communication	0.329	0.801
	Search	0.158	0.567
	AI Assistant	0.007	0.100
Top 100 Politicians	All	0.509	0.948
	News	0.021	0.079
	Social Media	0.192	0.640
	Browser	0.195	0.692
	Music and Video	0.197	0.726
	Communication	0.108	0.511
	Search	0.082	0.444
	AI Assistant	0.001	0.039
Top 100 Celebrities	All	0.857	0.990
	News	0.009	0.048
	Social Media	0.495	0.805
	Browser	0.283	0.806
	Music and Video	0.546	0.927
	Communication	0.233	0.758
	Search	0.138	0.566
	AI Assistant	0.009	0.066
Congressmembers	All	0.177	0.759
	News	0.004	0.037
	Social Media	0.042	0.366
	Browser	0.070	0.480
	Music and Video	0.035	0.355
	Communication	0.025	0.204
	Search	0.021	0.276
	AI Assistant	0.001	0.012

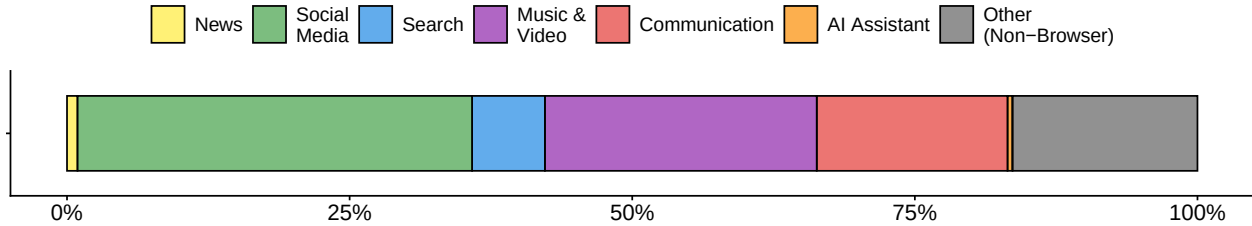
NOTES: The table reports the share of users exposed to each keyword set across app categories. “Average Day” refers to the mean daily share of users exposed on a typical day, while “Entire Period” denotes the share of individuals who encountered at least one term from the corresponding keyword set during the election campaign (September 1 – November 4, 2024). Keyword sets include primary political terms, the top 100 most observed politicians in our dataset, the top 100 most observed celebrities, and all U.S. congressmembers. For example, 0.40 under “Social Media” for the primary political terms indicates that 40% of individuals viewed election-related content on social media on an average day, while 0.99 under “All” for the top 100 celebrities shows that nearly all users encountered celebrity-related content at least once during the campaign period.

S3.3.1 Browser-Included News Visits

A key limitation in determining sources of information is that we cannot observe the websites individuals visit through their browsers, and these websites can belong to any category (e.g., facebook.com or nytimes.com). While it is likely that some exposure to news platforms occurs through browsers, based on previous research [61], most of the time spent browsing on mobile is spent on search (Google), music and video (YouTube), and various social media applications.

In this section, we provide additional estimates of the time spent on news platforms by following two procedures to impute news usage from browser time. The key assumption in each of these is that time spent on the browser is, for the most part, a distribution of time spent on applications on the phone and thus can be informative about what an individual is doing in the browser. The estimates in this section can be interpreted as expanding our measure of time spent on news apps to include time spent on news sites.

Figure S7: Share of Exposures by Category (Browser Imputation)



NOTES: The figure shows the overall distribution of election-related exposures across app categories during the period September 1–November 4, 2024, where we reallocate the exposures to browsers proportionally into other categories. The bar shows the average share of an individual’s total election-related exposures occurring in that category. The shares are computed in two steps: first, for each individual, exposures are summed by category/app and divided by the individual’s total election-related exposures; second, these individual-level shares are averaged across all devices. App categories include: News, Social Media, Search, Music and Video, Communication, AI Assistant, and Other (non-browser).

First, for the average individual, we use the time spent on all other application categories and make the assumption that the share of exposures across categories when people use their browser is the same as the overall share of exposures across the other categories. Supplementary Figure S7 presents the distribution of usage across categories under this assumption, which implies that news site usage only increases to 0.89% even accounting for browser time.

Second, we use a model-based imputation that exploits the fact that we not only observe the overall exposure levels across categories, but also the exact keywords across election-related and non-election-related content. In the rest of this section we describe a procedure leveraging that data to ascertain the fraction of news-site exposures in the browser.

There are $I = \{1, \dots, K\}$ non-Browser categories and $V = \{1, \dots, V\}$ set of keyword groups.

We define n_{iv} as the total counts of encounters in category i and keyword group v and, consequently, the empirical distribution of encounters across application categories as follows:

$$\phi_{iv} = \frac{n_{iv}}{\sum_{u=1}^V n_{iu}}, \quad \sum_{v=1}^V \phi_{iv} = 1.$$

where ϕ_{iv} is the share of encounters from keyword group v in category i for the average individual in the sample.

Our key assumption is that encounters on the browser are a convex combination of encounter distributions across application categories:

$$b_v(\beta) = \sum_{i=1}^K \beta_i \phi_{iv}, \quad \sum_{i=1}^K \beta_i = 1, \quad \beta_i \geq 0.$$

where $b_v(\beta)$ is the share of encounters from keyword group v in the browser category and β is the share of browser encounters generated from each of the $i \in I$ categories. We can write down the log likelihood based on the observed browser counts and the observed distribution of keyword encounters on non-browser applications as follows:

$$\ell(\beta) = \sum_{v=1}^V x_v \log\left(\sum_{i=1}^K \beta_i \phi_{iv}\right)$$

where x_v is the number of observed encounters across keyword group v within the browser for the average individual. We estimate β via maximum likelihood estimation using the expectation-maximization algorithm.

We know from Supplementary Section S7.2 that the density of exposures to any keywords (political and non-political) is higher on news than other categories. As such, to map our encounter estimates to time spent on different categories, we need to adjust for the fact that we will be more likely to capture keywords in some categories (such as news) as opposed to others. We define $\theta_i = \beta_i/m_i$ where m_i is the median fraction of time in category i that is captured by our full bank of keywords.²⁵ Therefore, θ_i provides us with a measure of the share of browser time spent on each of the categories. We estimate this on the data from September 1st until November 4th and restrict ourselves to $I = \{\text{News, Other}\}$ and estimate $\theta_{\text{News}} = 0.04$, meaning that approximately 4% of browser time is spent on news sites.

Our main interest, however, is understanding what fraction of election-related encounters on

²⁵We normalize θ_i such that $\sum_i \theta_i = 1$.

browsers comes from news platforms. We can compute this fraction using our estimates via the following formulation:

$$P(i = \text{News} \mid v = \text{election}) = \frac{\theta_{\text{News}} \phi_{\text{News, election}}}{\sum_{j \in I} \theta_j \phi_{j, \text{election}}}$$

In other words, we divide the time-adjusted share of news encounters in election-related encounters by the total time-adjusted election-related encounters. This provides us with an estimate of $P(i = \text{News} \mid v = \text{election}) = 0.21$, which tells us that a fraction of 0.21 of election-related exposures that we see in the browser category likely originate from news websites. Finally, we add these imputed news website exposures with the baseline 0.74% exposures in news applications and conclude that 4.93% of election-related exposures originate from either news applications or news websites visited in the browser.

To conclude, using both methods, we find that the mean share of exposures through news platforms increases from 0.74% to 0.89%-4.93%. Thus, this does not change our qualitative conclusion in the main text that consumption on news sites is low.

S4 Variance Decomposition

We conduct a variance decomposition in the spirit of [48]—hereafter, AKM.²⁶ We estimate the following empirical specification separately for each app category group g :

$$y_{iat}^g = \alpha_i^g + \beta_a^g + \gamma_t^g + \varepsilon_{iat}^g, \quad (\text{S1})$$

Our main outcome variable, y_{iat}^g , is the share of time user i was exposed to election-related content for app a (within group g) at time t .²⁷ This measure accounts for usage differences across individuals and applications—allowing us to isolate differences in the *composition* of the content that individuals are exposed to. The parameters α_i^g represent individual fixed effects, and can be interpreted as a combination of (time-invariant) portable individual characteristics that results in equal exposure to election-related content across all apps within the app group. The parameters β_a^g denote app fixed effects, which capture the possibility that some apps over- or under-expose their users systematically to election-related content. Lastly, γ_t^g denote time fixed effects, which allow election-related content to be popular at a given time period – for example, after presidential debates.

²⁶See [62] and [63] for other media applications of the AKM decomposition.

²⁷Supplementary Section S7.2 explains how we use our set of keywords to calculate the total time spent on each app, the denominator of our outcome variable.

We interpret this estimation in descriptive terms, not as a causal analysis. For a causal interpretation, this methodology imposes three main assumptions [64]. First, the error terms ε_{iat}^g are mean-independent from the individual, app, and time-fixed effects. This “exogenous mobility” assumption allows for unrestricted and arbitrarily correlated sorting patterns into apps based on individual and app fixed effects, but rules out sorting based on the error term. Second, an additive separability assumption rules out interactions between individual and app fixed-effects [65]. This assumption does not rule out the possibility of content personalization. It rather rules out *differential* personalization, such as individuals with a strong preference for election content experiencing relatively more exposures on higher-exposure apps than in lower-exposure apps, within the same category. Third, following [66], app effects are identified only within a “connected set” of apps linked by individual “moves”. In our setting, this is not an issue as each individual uses many apps. In other words, multi-homing mitigates limited mobility bias [67]. Nevertheless, to ensure reliable estimates, we restrict the sample to individuals visiting at least 30 apps and to apps visited by at least 30 such individuals.

As in AKM, we decompose the variance of election-related exposures using Equation (S1), within each app group g :

$$\begin{aligned}\text{Var}(y_{iat}^g) = & \text{Var}(\alpha_i^g) + \text{Var}(\beta_a^g) + \text{Var}(\gamma_t^g) + \text{Var}(\varepsilon_{iat}^g) \\ & + 2\text{Cov}(\alpha_i^g, \beta_a^g) + 2\text{Cov}(\alpha_i^g, \gamma_t^g) + 2\text{Cov}(\beta_a^g, \gamma_t^g).\end{aligned}\quad (\text{S2})$$

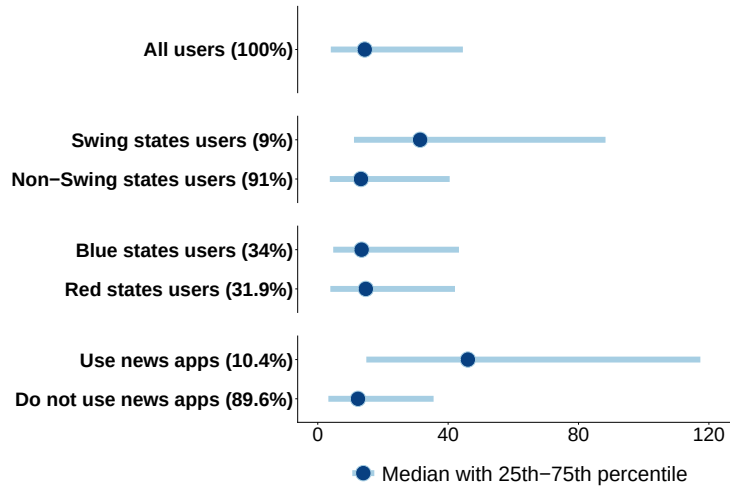
By comparing $\text{Var}(\alpha_i^g)$ and $\text{Var}(\beta_a^g)$, we assess whether individual characteristics or systematic app-specific factors play a larger role in shaping content exposure. The covariance term $\text{Cov}(\alpha_i^g, \beta_a^g)$ indicates the pattern of individual sorting into apps. For example, if this component is positive, it reflects that individuals with particularly strong exposure to election-related content also tend to use apps that provide frequent coverage of such content.²⁸ We estimate the variance decomposition within each app group g . To aggregate across app groups, we weight the share of exposures of each group by the time spent on that group.

²⁸We conducted a simulation exercise that assumes that the true data-generating process is either entirely determined by (a) applications or (b) individuals and confirm that our procedure correctly identifies which effects dominate.

S5 Additional Results for Drivers of Consumption

S5.1 Heterogeneity Across Individuals

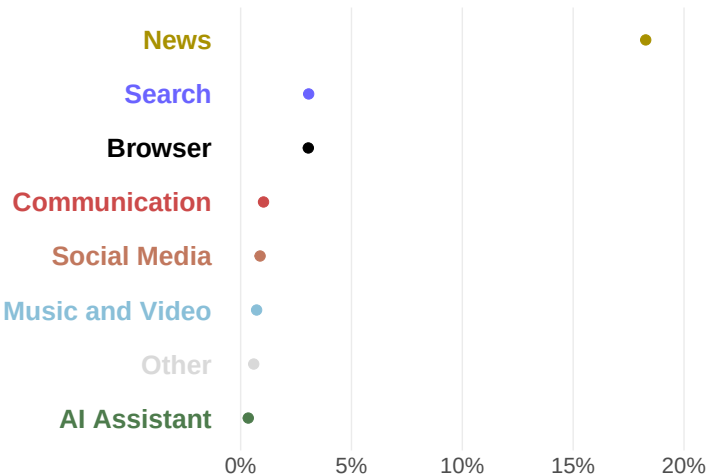
Figure S8: Individual Heterogeneity in Election-Related Encounters



NOTES: The figure shows the distribution of average daily encounters of the median individual with election-related content across subgroups (September 1st–November 4th 2024). Dots represent group medians; bars indicate interquartile ranges (25th–75th percentile). Subgroup labels include the group's share of the total sample in parentheses. Geographic groups were assigned based on the overrepresented state method in encounter data (see Supplementary Section [S7.1](#)).

S5.2 Heterogeneity Across Applications

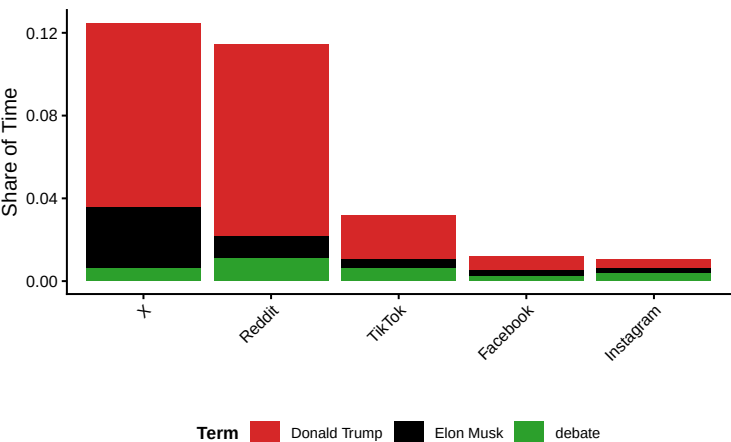
Figure S9: Heterogeneity in Time Spent on Election-Related Content across Categories



NOTES: The figure illustrates variation in the share of time users spent on election-related content across app categories. The share is computed as estimated time spent observing election-related content based on exposures divided by estimated time usage on the application. The time usage is based on exposures to the full set of keywords and the imputation in Supplementary Section S7.2. The figure reflects aggregate user behavior between September 1 and November 4, 2024.

S5.3 Additional Variance Decomposition Results

Figure S10: Elon Musk and Donald Trump Share



NOTES: This figure presents the share of total encounters, aggregated across individuals, of the terms 'Donald Trump', 'Elon Musk', and 'debate' across each of the presented applications. The figure uses data from September 1st until November 4th.

Table S3: Variance decomposition (Weekly, election-related keyword encounters)

	All apps	Social	Communication
Individual FE	0.86	0.76	0.81
App FE	0.17	0.24	0.13
Time FE	0.04	0.02	0.02
Covariance of ind. & app FE	-0.15	-0.00	0.05
Covariance of ind. & time FE	0.02	-0.02	-0.01
Covariance of app & time FE	0.07	0.00	-0.00
Adj. R^2	0.24	0.26	0.22
N apps	285	10	13
N individuals	1469	1469	1469

NOTES: This table presents the share of explained variance corresponding to each component of the variance decomposition equation. We estimate this equation separately for each app category. The “All apps” column aggregates across all app categories: AI Assistant, Browser, Communication, Music & Video, News, Search, and Social Media. This aggregation is done by weighting the share of encounters of each group. We exclude apps with fewer than 30 unique users and individuals with fewer than 30 unique apps.

Table S4: Variance decomposition (Daily, election-related keyword exposures)

	All apps	Social	Communication
Individual FE	0.81	0.71	0.83
App FE	0.17	0.24	0.10
Time FE	0.03	0.02	0.02
Covariance of ind. & app FE	-0.06	0.04	0.07
Covariance of ind. & time FE	-0.01	-0.01	-0.01
Covariance of app & time FE	0.07	0.00	-0.00
Adj. R^2	0.24	0.26	0.21
N apps	285	10	13
N individuals	1469	1469	1469

NOTES: This table presents the share of explained variance with observations at the daily level instead of weekly level. We estimate this equation separately for each app category. The “All apps” column aggregates across all app categories. We exclude apps with fewer than 30 unique users and individuals with fewer than 30 unique apps.

Table S5: Variance decomposition (Aggregated across time, election-related keyword exposures)

	All apps	Social	Communication
Individual FE	0.87	0.74	0.81
App FE	0.24	0.27	0.14
Time FE	0.00	0.00	0.00
Covariance of ind. & app FE	-0.11	-0.02	0.05
Covariance of ind. & time FE	0.00	0.00	0.00
Covariance of app & time FE	0.00	0.00	0.00
Adj. R^2	0.12	0.23	0.17
N apps	285	10	13
N individuals	1469	1469	1469

NOTES: This table presents the share of explained variance without time fixed effects. We estimate this equation separately for each app category. The “All apps” column aggregates across all app categories. We exclude apps with fewer than 30 unique users and individuals with fewer than 30 unique apps.

Table S6: Variance decomposition (Weekly, election-related keyword exposures, all apps)

	All apps	Social	Communication
Individual FE	0.70	0.65	0.80
App FE	0.19	0.31	0.13
Time FE	0.02	0.02	0.03
Covariance of ind. & app FE	0.00	0.03	0.06
Covariance of ind. & time FE	0.03	-0.01	-0.01
Covariance of app & time FE	0.06	0.00	-0.00
Adj. R^2	0.25	0.26	0.20
N apps	331	10	14
N individuals	2245	2245	2245

NOTES: This table presents the share of explained variance without restricting to apps with ≥ 30 unique users. We exclude individuals with fewer than 30 unique apps.

Table S7: Variance decomposition (election-related exposures, [49] correction)

	Social	Communication
Individual FE	0.56	0.65
App FE	0.47	0.25
Covariance of ind. & app FE	-0.03	0.10
Adj. R^2	0.29	0.22
N apps	10	13
N individuals	1469	1469

NOTES: This table presents the share of explained variance using the leave-one-out bias correction of [49]. We rely on the implementation from <https://github.com/vahid-moghani89/LeaveOutKSS-R>. We consider this for election-related exposures and aggregate across time. We only report the category-specific estimates since the library implementation does not compute the covariance terms needed to do our category weighting procedure. We exclude apps with fewer than 30 unique users and individuals with fewer than 30 unique apps.

S5.4 Content Consumption Survey

This section provides additional details on the survey described in the main text.

Our survey was completed by 1,011 respondents recruited in August 2025 through Cloud Research, a popular online panel provider [68]. We targeted adults in the United States on their mobile phone and respondents who report having an Android device in order to match as closely as possible our primary sample. We filter our final sample to those who passed our three attention checks, including a hidden question for detection of bots, which leaves 986 respondents.

The first block of interest asked participants whether they felt they consumed the right amount of different types of content on various apps, or whether they over- or under-consumed them. We asked participants about content related to politics and Elon Musk, as well as other topics (such as entertainment) and people (such as Mark Zuckerberg). Concretely, we asked:

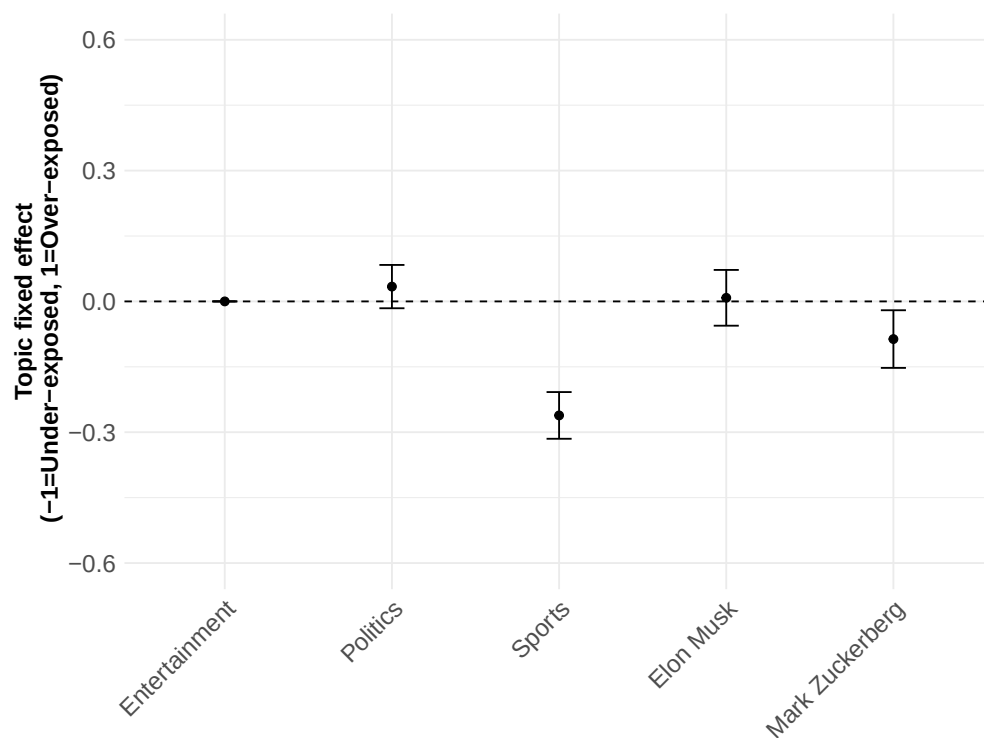
Let us know about how you feel about the amount of content related to [topic/person] you receive in your feed on each platform. If you have stopped using any of these platforms, please consider the last time you used them within the past year.

We filtered out responses for respondents who stated that they did not use a given application.

Our empirical analysis relies on estimating the following equation reported in the main text. By controlling for respondent fixed effects, we isolate respondents' sense of exposure to specific content in each app after taking into account their overall tendency to feel over- or under-exposed across all applications. This is similar to the regressions in the main text that isolate differences in content diets across applications. However, while in the main text the differences across applications could be driven by differences in how individuals use the applications, here the differences

are in the discrepancy between desired and actual content exposure. Thus, the estimated coefficient captures whether, for a given application, respondents are more exposed to a type of content than they would like to be. A positive estimate for a given application indicates perceived over-exposure, relative to the perceived over-exposure levels across all applications. If an application has a non-zero coefficient systematically across all topics, it is possible that the application is difficult to personalize. However, if an application has an estimate of zero, relative to other applications, for most topics and only a positive estimate for some, then it is possible that it over-prioritizes these types of content.

Figure S11: Overall Exposure to Content



NOTES: This figure presents topic fixed effects from the overall exposure regression. Entertainment is the omitted category, serving as the reference group. Positive values indicate higher self-reported over-exposure relative to entertainment; negative values indicate under-exposure. The 95% confidence intervals are constructed using standard errors clustered at the respondent level.

The second block of interest asked respondents whether, across all applications on their phone, they felt they consumed too little (-1), the right amount (0), or too much (1) of a particular topic on their phone. The resulting estimates are plotted in Figure S11. Respondents report similar amounts of over-consumption of politics as they do of entertainment when considering all applications together, in spite of us finding self-reported over-consumption of politics when considering social media applications individually. These results are consistent with our finding that individual fixed

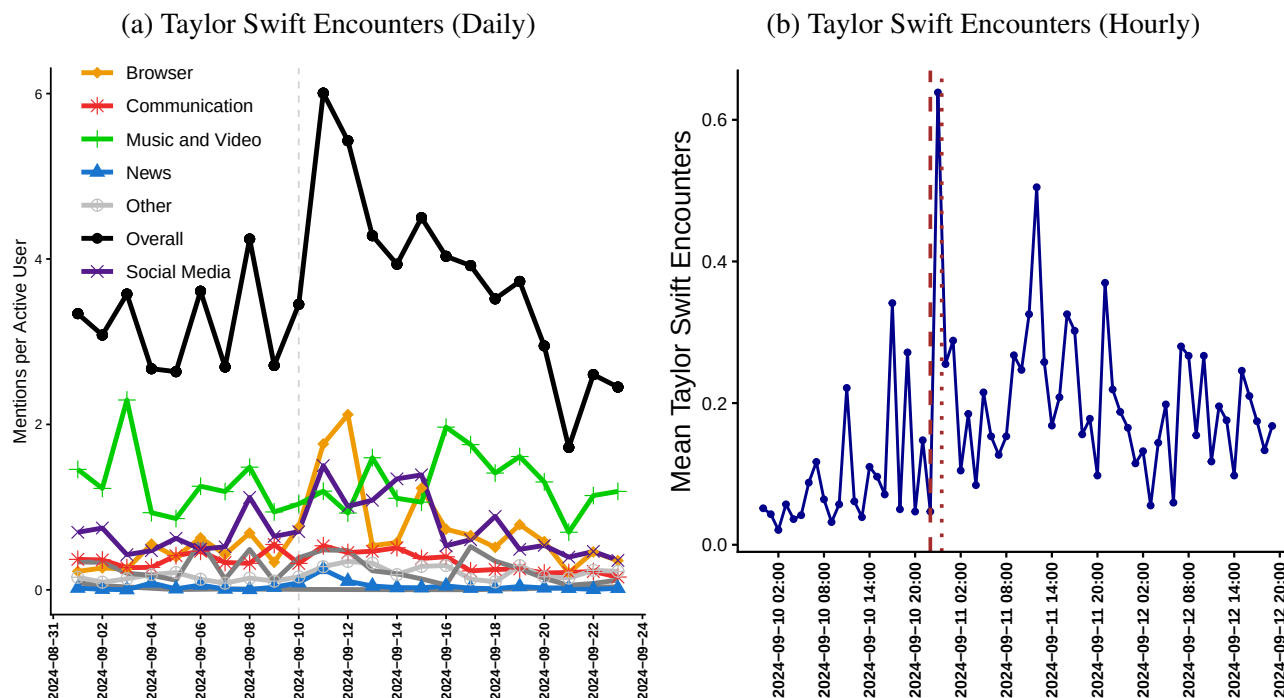
effects are the main driver of heterogeneity in consumption, as the content individuals receive in aggregate tends to match their preferences.

Putting these two results together, they paint a similar picture as the main text, and suggest possible mechanisms. Given that 1) individuals' satisfaction with the amount of political content varies widely across apps, and 2) many are dissatisfied with their exposure (65% report over-consuming Elon Musk-related content on X), our findings suggest that not all platforms are delivering the content users want, and that over-consumption of political and Elon Musk-related content on X is at least partly driven by application prioritization choices. At the same time, similarly to our variance decomposition exercise, the magnitudes of consumption aggregated across different applications suggest that, despite these systematic application effects, individuals are still able to curate their full set of applications to produce desired consumption levels of political content.

The third block of the survey additionally contains questions regarding screen management software and self-reported levels of news consumption, with the analysis of this in Supplementary Section [S2](#).

S6 Taylor Swift Endorsement Analysis

Figure S12: Taylor Swift Encounters



NOTES: Subfigures (a) and (b) describe the average daily and hourly encounters, respectively, with the term “Taylor Swift” per active individual across different app categories. The dashed vertical line in the first panel marks the timing of Swift’s endorsement, while in the second panel the dashed vertical marks the start of debate and the dotted line marks Swift’s endorsement.

In this section, we analyze the effect of Taylor Swift’s Instagram endorsement of Kamala Harris. The endorsement occurred directly following the highly anticipated presidential debate between Kamala Harris and Donald Trump, which began on September 10th, 2024 at 9:00 PM ET. Taylor Swift’s endorsement occurred minutes after the conclusion of the debate at 10:30 PM ET when she posted her endorsement of Kamala Harris on Instagram and directed her followers to register to vote.^{29,30}

To study whether Swift’s endorsement could have affected election-related content consumption, we begin by documenting the number of mentions for the term ‘Taylor Swift’, highlighting

²⁹For the endorsement itself, see https://www.instagram.com/p/C_wtAOKOW1z/?hl=en. For a discussion of the endorsement, see <https://www.nytimes.com/2024/09/10/us/taylor-swift-endorses-kamala-harris.html>.

³⁰Reportedly, the number of visits to ‘vote.gov’ increased dramatically following the debate and the endorsement. See <https://www.cbsnews.com/news/taylor-swift-kamala-harris-endorsement-vote-gov/>.

that individuals in our sample were exposed to the endorsement event. Figure S12a shows that overall mentions of Taylor Swift were fairly flat from September 1st until September 10th, while the number of mentions for Taylor Swift dramatically spiked on September 11th, increasing by over 100%. As the endorsement was made late in the evening on September 10th, this is when we expect to see that most individuals would initially be exposed to it. Figure S12b shows the number of encounters with the term ‘Taylor Swift’ by hour and confirms the interest spiked exactly when Swift made her endorsement, but was persistently higher for several days afterward. Indeed, the overall mentions of Taylor Swift do not return back to pre-endorsement levels until September 20th, nearly 10 days after the endorsement, suggesting that the event could also have a persistent effect on election-related content consumption. While the spikes for Taylor Swift on news applications only persist for the day following the endorsement, the breakdown by category suggests longer-lasting elevated exposures to Taylor Swift on social media and communication applications.

In order to measure the effect of the endorsement on election-related content consumption, we use a difference-in-differences analysis comparing individuals who were relatively more likely to interact with Taylor Swift-related content. We define ‘Swifties’ as individuals in our dataset whose share of ‘non-political’ exposures to the term ‘Taylor Swift’ (i.e., filtering out instances with a political surrounding word) is above the median.³¹ We consider the following regression specification to assess the overall causal effect of the endorsement for a given individual i and time period t :

$$Y_{i,t} = \sum_t \beta_t (\text{Day}_t \times S_i) + \gamma_t + \kappa_i + \epsilon_{it} \quad (\text{S3})$$

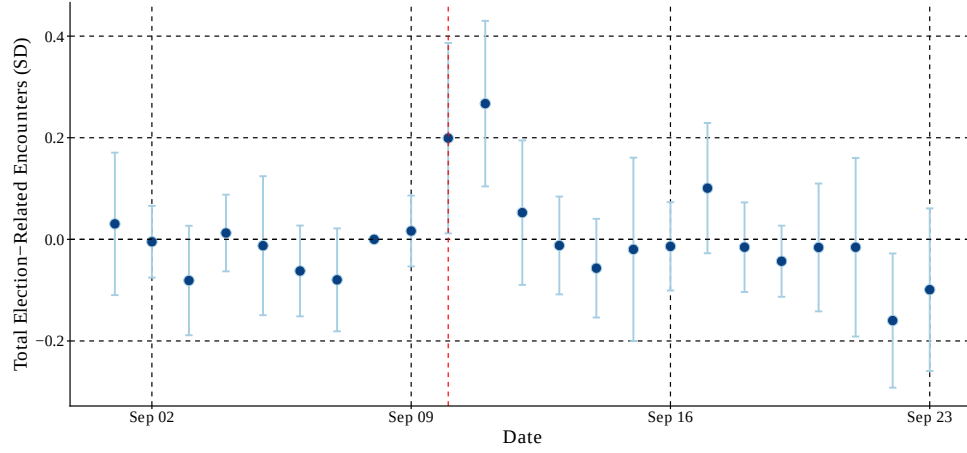
where S_i indicates that individual i is a Swiftie (treated), γ_t denotes daily fixed effects, and κ_i denotes individual fixed effects. We cluster standard errors at the individual level.

Figure S13 plots the estimated treatment effects, β_t , derived from estimating specification (S3) with normalized election-related encounters and exposures as Y_{it} . Figures S13a and S13b show that there was a statistically and economically significant increase in both encounters and exposures, respectively, the two days following the debate. Notably, there is a 0.27 and 0.33 standard deviation increase in election-related encounters and exposures on September 11th, respectively, the day following the debate. The positive effects on election-related content consumption persist for several days and dissipate over time, eventually returning to pre-endorsement levels.

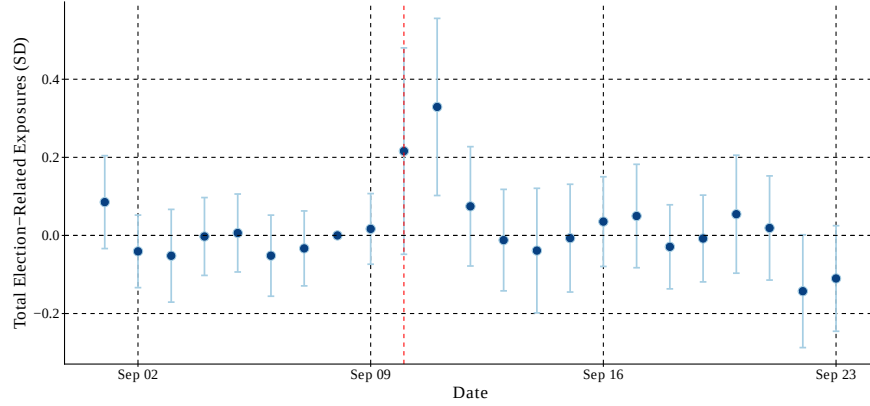
³¹We compute the share for each individual using their total exposures to Taylor Swift divided by the total number of exposures across all terms. Our results are robust to considering an alternative definition of ‘Swifties’ as individuals with at least one non-political encounter with the term ‘Taylor Swift’ in a music & video applications such as Spotify and YouTube. Given the popularity of Taylor Swift, this leads to 30.5% of the individuals that we observe being classified as ‘Swifties’.

Figure S13: Effect of Swift Endorsement on Election-Related Content Consumption

(a) Election-Related Encounters



(b) Election-Related Exposures



NOTES: This figure plots the estimated daily effect of Taylor Swift's public endorsement of Kamala Harris on election-related encounters (panel a) and exposures (panel b), using estimates from specification (S3). Each point represents the estimated difference in standardized daily election-related exposures or encounters between individuals with above- and below-median pre-endorsement exposure share to the term 'Taylor Swift'. The red dashed line marks the date of the endorsement (September 10, 2024). 95% confidence intervals of the estimated treatment effect are shown, derived from standard errors clustered at the individual level.

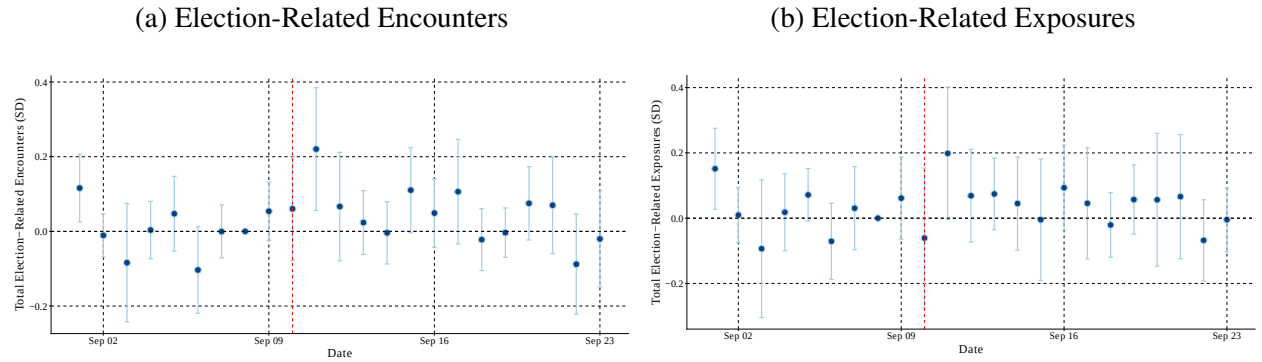
We consider an additional robustness specification that controls for the fact that more politically active individuals may differentially respond to the endorsement. We estimate the following regression specification:

$$Y_{i,t} = \sum_t \beta_t (\text{Day}_t \times S_i) + \alpha_t (\text{Day}_t \times Y_{i,-1}) + \kappa_i + \epsilon_{it} \quad (\text{S4})$$

where $Y_{i,-1}$ is the mean value of the outcome variable in the pre-period, which we interact with the date in order to control for differential responsiveness to unrelated time-varying shocks that can influence election-related content consumption. In particular, we proxy and control for political activeness using election-related content consumption in the baseline and interact this with the day fixed effects. Figure S14 plots the estimated treatment effects, β_t , derived from estimating specification (S4) and find similar patterns as our primary specification.

While the results in the main text indicate that there are few events that increase election-related content consumption, our results here indicate that endorsements, and more broadly exposure to election-related content via non-political figures, can lead to increases in election-related content consumption among specific segments of the population.

Figure S14: Effect of Swift Endorsement on Election-Related Content Consumption (Robustness)



NOTES: This figure plots the estimated daily effect of Taylor Swift's public endorsement of Kamala Harris on election-related encounters (panel a) and exposures (panel b), using estimates from specification (S4). Each point represents the estimated difference in standardized daily election-related exposures (encounters) between individuals with above- and below-median pre-endorsement exposure (encounters) share to the term 'Taylor Swift'. The red dashed line marks the date of the endorsement (September 10, 2024). 95% confidence intervals of the estimated treatment effect are shown, derived from standard errors clustered at the individual level.

S7 Data Imputation Procedures

In this section we provide details for two imputation procedures that we use throughout the text: imputation of the state of an individual and the total time spent on the phone and individual applications.

S7.1 State Classification

In this section we discuss our procedure for classifying the state of residence of each individual in our sample. In order to do so, we make use of two aspects of the data: the full set of states as

keyword encounters and the timezone set on each individual’s device.

To assign each user a likely U.S. state of residence, we implemented a classification procedure based on the overrepresentation of U.S. state mentions in our encounter data. We first identified encounters involving state terms and for each device, we computed the relative share of mentions for each state term and compared it to the corresponding global share of that state’s mentions across all devices. This ratio, interpreted as an overrepresentation score, captures how much more frequently a user mentions a specific state compared to the population baseline. We then selected the top three overrepresented states for each device, reflecting the strongest state-level signal based on their appearance on the user’s screen.

To improve classification accuracy, we leveraged time zone data provided by Screenlake for each device and cross-referenced it with the local time zones associated with each U.S. state.³² For example, California and Oregon were mapped to the Pacific Time Zone, while Texas and Illinois corresponded to Central Time. The classification method evaluated the top three overrepresented states for each device in descending order. It first attempted to assign the top-ranked state (i.e., the one most disproportionately mentioned for the user), and verified whether its assigned time zone matched the device’s time zone. If the first choice was inconsistent with the device’s time zone, we evaluated the second most overrepresented state, and if necessary, the third. If none of the top three states aligned with the device’s time zone, the algorithm defaulted to the top-ranked state. This matching process produced a proposed state classification for each device.

Based on this proposed state variable, we assign each individual to one of three political categories: swing states, blue states, and red states.

1. Swing states: Arizona, Georgia, Michigan, Nevada, Pennsylvania, and Wisconsin.
2. Blue states: California, Colorado, Connecticut, Delaware, Hawaii, Illinois, Maine, Maryland, Massachusetts, Minnesota, New Hampshire, New Mexico, New York, Oregon, Rhode Island, Vermont, and Washington.³³
3. Red states: Alabama, Alaska, Arkansas, Florida, Idaho, Indiana, Iowa, Kansas, Kentucky, Louisiana, Mississippi, Missouri, Montana, Nebraska, North Dakota, Ohio, Oklahoma, South Dakota, Tennessee, Texas, Utah, and Wyoming.

Several states were excluded from the final political classification due to ambiguity in keyword identification. Specifically, references to “Carolina” and “Virginia” could not be reliably disaggregated into North vs. South Carolina or Virginia vs. West Virginia, leading to their exclusion from the swing/blue/red categorization to avoid misclassification. By contrast, the term “Dakota”

³²Screenlake infers time zone from the various brands, apps, and/or terms that appear on the screen.

³³Data on New Jersey was not collected due to a technical error.

was not problematic for classification purposes, as both North and South Dakota are consistently categorized as red states.

Our final validation step suggests that this classification method correctly matches the user's time zone in 86% of cases.

S7.2 Time Usage Imputation from Encounters Data

Table S8: Fraction of Time Captured by Exposures

Aggregation	Mean	Min	25th	Median	75th	Max
Total Phone Time	0.0407	0.0012	0.0219	0.0350	0.0528	0.4615
Apps						
Facebook	0.0782	0.0005	0.0246	0.0465	0.0840	0.9730
Gmail	0.1872	0.0016	0.0748	0.1383	0.2418	1.0000
Google	0.1860	0.0013	0.0853	0.1449	0.2384	1.0000
Instagram	0.0364	0.0005	0.0162	0.0264	0.0413	0.5593
Messages	0.1081	0.0010	0.0289	0.0572	0.1121	1.0000
NYTimes	0.1402	0.0032	0.0606	0.0722	0.1288	0.6250
Reddit	0.1001	0.0009	0.0312	0.0587	0.1098	1.0000
Spotify	0.1449	0.0002	0.0498	0.1010	0.1913	1.0000
Threads	0.1363	0.0054	0.0347	0.0887	0.1793	0.9474
TikTok	0.0486	0.0005	0.0192	0.0329	0.0549	0.7178
WhatsApp	0.0486	0.0008	0.0152	0.0305	0.0522	0.8308
X	0.0528	0.0007	0.0110	0.0210	0.0499	0.7500
YouTube	0.0362	0.0002	0.0103	0.0198	0.0393	0.9375
Categories						
AI Assistant	0.1190	0.0009	0.0312	0.0789	0.1558	0.8165
Browser	0.0930	0.0003	0.0328	0.0671	0.1154	1.0000
Communication	0.0509	0.0006	0.0178	0.0352	0.0612	0.8571
Music and Video	0.0501	0.0002	0.0125	0.0253	0.0532	1.0000
News	0.2508	0.0032	0.0736	0.1738	0.3142	1.0000
Other	0.0658	0.0002	0.0237	0.0466	0.0862	1.0000
Search	0.1835	0.0013	0.0824	0.1395	0.2329	1.0000
Social Media	0.0422	0.0003	0.0181	0.0313	0.0487	1.0000

NOTES: This figure presents the fraction of time captured by exposures, using the joined encounter and time usage data between December 19th, 2024 until January 25th, 2025. The first row computes this for overall phone usage, the next set of rows for individual applications, and the final set of rows for application categories.

In this section we describe how we use the encounters data to provide an approximation of the time spent on different applications, application categories, and on the phone in general. The encounters and time usage streams are separate data streams that Screenlake collects from the phone. The

permissions that users are required to enable are different for each stream. Unfortunately, during the election campaign, we do not have reliable time usage data. Therefore, for this imputation, we use data from December 19th, 2024 until January 25th, 2025 when we have reliable time usage data and can accurately match it with the encounters data.

We use the following imputation procedure. First, as in the main text, we approximate the time spent using the encounters data by multiplying the number of app-exposure-date observations by three, since the app records the on-screen encounters every three seconds. Then, we compute for each individual their total phone time in which they were exposed to any keyword (based on our full list of thousands of keywords and not just the subset of election-related keywords) according to the unique number of exposures in a given day. We drop the individuals for whom this results in a time estimate higher than the application usage data, since it is likely that these individuals consistently have enabled the encounter permissions and not the time usage permissions. Fortunately, only 44 individuals are dropped based on this criterion. For each of the remaining 1,980 individuals, we aggregate the time spent on each app over the entire time period based on the application data, and aggregate the time spent on each app while exposed to any keyword based on the encounter data. We then compute the fraction of time captured by the full set of keyword encounters overall on the phone, by top applications, and by categories. We use the median of these estimates throughout the text when we approximate time spent using our encounter data. The results are presented in Table S8.

S8 Expert Survey

To assess perceptions among experts, we conducted an online survey on opinions about political consumption between November 6, 2025, and December 17, 2025. The survey targeted researchers with expertise in political economy, political behavior, and related fields, as well as experienced forecasters. Our recruitment strategy combined three complementary channels to ensure a diverse pool of respondents with varying backgrounds and perspectives.

First, we constructed a sample of experts in political economy. We scraped all papers listed under the Political Economy program of the NBER Working Papers series from January 2020 through October 2025, prior to the launch of our survey. From this list, we extracted all coauthors and emailed a random subsample of these researchers with invitations to participate in the survey from November 7 to December 3, 2025.

Second, we listed the study on the Social Science Prediction Platform, a commonly used platform with a pool of individuals who regularly participate in forecasting exercises of social science research (see, for example, Enke et al. 69, Deshpande and Dizon-Ross 70, DellaVigna et al. 71). This channel allowed us to reach experienced forecasters who may not specialize in the substan-

tive topic of the paper but who have extensive experience forming and evaluating probabilistic predictions.

Third, a link to the survey was distributed in an email via the POLMETH mailing list on December 4, 2025. The POLMETH listserv reaches a broad community of scholars in political methodology and related areas, including faculty, postdoctoral researchers, and advanced graduate students. The topic of the survey was not disclosed in the email in order to reduce selection into the sample.

In total, we received 279 completed responses. Of these, 36 respondents were recruited through the NBER-based email outreach, 180 respondents through the POLMETH mailing list, and 63 respondents through SSPP.

We report results for the main prediction questions, excluding participants who reported prior exposure to the paper or gave internally inconsistent answers.³⁴ We find that expert predictions diverge substantially from the actual data. For instance, while the actual median exposure to election-related content is less than one minute, the median expert prediction is 45 minutes. Predictions for relative levels are similarly miscalibrated: experts tend to overestimate both the share of content related to politics and the share of election-related exposure coming from news sources. In terms of the celebrity-to-politician ratio, the median individual in our data was exposed to 9.6 times more sports and entertainment celebrities than politicians, yet the median expert prediction for this ratio is only 3.5 – implying that experts perceive a more even balance between celebrity and political content than actually exists.

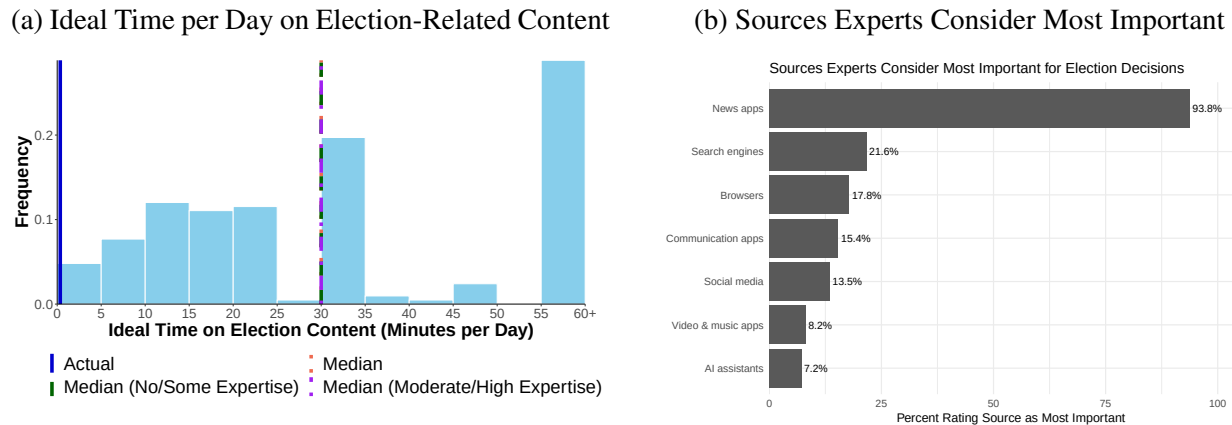
We also find substantial heterogeneity in forecast accuracy across experts. Using responses to the question, “How would you describe your expertise in media consumption?”, we find that experts with greater self-reported expertise produce more accurate predictions for the level of exposure and ratio of celebrities to politicians measures. Those who report having conducted related research or actively publish in media consumption predict lower exposure, and a higher ratio of celebrities to politicians. Interestingly, the median prediction for the share of election-related exposure coming from news sites or apps is 15% across expertise levels.

Taken together, the expert survey suggests that the low levels of political exposure we document, and the dominance of entertainment content and non-news app sources, are new findings rather than conventional wisdom.

Beyond predictions about actual consumption, the expert survey also reveals a stark gap between observed behavior and what experts consider ideal. Figure S15 summarizes expert views on how people should consume election-related news to make an informed decision. The median

³⁴Specifically, we dropped 5 participants who reported having “read part of the paper or seen it presented” and 51 participants who predicted higher exposure to presidential candidates than to election-related content overall, which may indicate a misunderstanding of the question.

Figure S15: Expert Views on Ideal News Consumption for Informed Election Decisions



Notes: Panel (a) shows the distribution of expert responses to the question of how many minutes per day people should have spent on election-related content between September 1 and November 4, 2024 to make informed decisions about the 2024 presidential election. Panel (b) shows the percentage of experts who rated each source as most important for making an informed election decision; respondents who rated multiple sources equally highly are counted once for each tied source, so percentages sum to more than 100%.

expert believes people should spend around 30 minutes per day on election-related content (Panel a), substantially more than the level of actual exposure we document. Experts also overwhelmingly agree on the ideal source: 93.8% rate news apps as the most important source for making an informed election decision, far exceeding any other category (Panel b). Notably, these normative views are remarkably consistent across expertise levels, in contrast to the heterogeneity we document for predictions about actual consumption. The contrast between these normative views and the descriptive patterns documented earlier—low overall exposure, dominated by entertainment content delivered through non-news platforms—underscores the magnitude of the gap between how experts believe an informed electorate should engage with election news and how Americans actually do.

S9 Dictionaries

S9.1 Lists of Detected Keywords

In this section, we provide additional documentation of the full set of election-related keywords used in our analyses. Outside of election-related keywords, Screenlake captures a large number of nonpolitical keywords, which we use as a proxy for time spent on the phone (discussed in Section S7.2). In total, Screenlake captures 2286 keywords across sports (e.g., Pittsburgh Steelers), entertainment (e.g., NCIS), celebrities (e.g., Gal Gadot), brands (e.g., Sprite), countries (e.g., Israel), and miscellaneous (e.g., weather).

The full set of political figures that Screenlake captures and we include in our definition of election-related keywords are as follows:

- **Presidential Candidates:** Donald Trump, Kamala Harris, J.D. Vance, Tim Walz
- **Notable Political Figures:** Bernard Sanders, Hillary Clinton, Joe Biden, Nancy Pelosi, Nikki Haley, Robert F Kennedy Jr
- **Governors:** Albert Bryan, Andy Beshear, Arnold Palacios, Bill Lee, Brad Little, Brian Kemp, Chris Sununu, Dan McKee, Doug Burgum, Eric Holcomb, Gavin Newsom, Glenn Youngkin, Greg Abbott, Greg Gianforte, Gretchen Whitmer, Henry McMaster, Janet Mills, Jared Polis, Jay Inslee, JB Pritzker, Jim Justice, Jim Pillen, Joe Lombardo, John Carney, Josh Green, Josh Shapiro, Kathy Hochul, Katie Hobbs, Kay Ivey, Kevin Stitt, Kim Reynolds, Kristi Noem, Laura Kelly, Lemanu Peleti, Lou Leon, Mark Gordon, Maura Healey, Michelle Lujan, Mike DeWine, Mike Dunleavy, Mike Parson, Ned Lamont, Pedro Pierluisi, Phil Murphy, Phil Scott, Ron DeSantis, Roy Cooper, Sarah Huckabee, Spencer Cox, Tate Reeves, Tina Kotek, Tony Evers, Wes Moore
- **Congressmembers:** Aaron Bean, Abigail Spanberger, Adam Schiff, Adam Smith, Adrian Smith, Adriano Espaillat, Al Green, Alex Padilla, Alexandria Ocasio-Cortez, Alma Adams, Ami Bera, Amy Klobuchar, André Carson, Andrea Salinas, Andrew Clyde, Andrew Garbarino, Andy Barr, Andy Biggs, Andy Harris, Andy Kim, Andy Ogles, Angie Craig, Angus S. King, Ann Wagner, Anna Eshoo, Anna Paulina Luna, Annie Kuster, Anthony D’Esposito, Ashley Hinson, August Pfluger, Austin Scott, Ayanna Pressley, Barbara Lee, Barry Loudermilk, Barry Moore, Becca Balint, Ben Cline, Ben Ray Luján, Benjamin L. Cardin, Bennie Thompson, Bernard Sanders, Beth Van Duyne, Betty McCollum, Bill Cassidy, Bill Foster, Bill Hagerty, Bill Huizenga, Bill Keating, Bill Pascrell, Bill Posey, Blaine Luetkemeyer, Blake Moore, Bob Good, Bob Latta, Bobby Scott, Bonnie Watson Coleman, Brad Finstad, Brad Schneider, Brad Sherman, Brad Wenstrup, Brandon Williams, Brendan Boyle, Brett Guthrie, Brian Babin, Brian Fitzpatrick, Brian Mast, Brian Schatz, Brittany Pettersen, Bruce Westerman, Bryan Steil, Buddy Carter, Burgess Owens, Byron Donalds, Carlos A. Giménez, Carlos Giménez, Carol Miller, Catherine Cortez Masto, Cathy McMorris Rodgers, Celeste Maloy, Charles E. Schumer, Chip Roy, Chris Deluzio, Chris Pappas, Chris Smith, Chris Van Hollen, Chrissy Houlahan, Christopher A. Coons, Christopher Murphy, Chuck Edwards, Chuck Fleischmann, Chuck Grassley, Chuy García, Cindy Hyde-Smith, Claudia Tenney, Clay Higgins, Cliff Bentz, Colin Allred, Cori Bush, Cory A. Booker, Cory Mills, Cynthia M. Lummis, Dale Strong, Dan Bishop, Dan Crenshaw, Dan Goldman, Dan Kildee, Dan

Meuser, Dan Newhouse, Dan Sullivan, Daniel Webster, Danny Davis, Darin LaHood, Darrell Issa, Darren Soto, David Joyce, David Kustoff, David Rouzer, David Schweikert, David Scott, David Trone, David Valadao, Dean Phillips, Deb Fischer, Debbie Dingell, Debbie Lesko, Debbie Stabenow, Debbie Wasserman Schultz, Deborah Ross, Delia Ramirez, Derek Kilmer, Derrick Van Orden, Diana DeGette, Diana Harshbarger, Dina Titus, Don Bacon, Don Beyer, Don Davis, Donald Norcross, Doris Matsui, Doug LaMalfa, Doug Lamborn, Drew Ferguson, Dusty Johnson, Dutch Ruppersberger, Dwight Evans, Ed Case, Edward J. Markey, Eli Crane, Elise Stefanik, Elissa Slotkin, Elizabeth Warren, Emanuel Cleaver, Emilia Sykes, Eric Burlison, Eric Schmitt, Eric Sorensen, Eric Swalwell, Erin Houchin, Frank Lucas, Frank Pallone, Frederica Wilson, French Hill, Gabe Vasquez, Garret Graves, Gary C. Peters, Gary Palmer, Gerry Connolly, Glenn Grothman, Glenn Ivey, Glenn Thompson, Grace Meng, Grace Napolitano, Greg Casar, Greg Landsman, Greg Lopez, Greg Murphy, Greg Pence, Greg Stanton, Greg Steube, Gregory Meeks, Gus Bilirakis, Guy Reschenthaler, Gwen Moore, Hakeem Jeffries, Hal Rogers, Haley Stevens, Hank Johnson, Harriet Hageman, Henry Cuellar, Hillary Scholten, Ilhan Omar, Jack Bergman, Jack Reed, Jacky Rosen, Jahana Hayes, Jake Ellzey, Jamaal Bowman, James Comer, James E. Risch, James Lankford, Jan Schakowsky, Jared Golden, Jared Huffman, Jared Moskowitz, Jasmine Crockett, Jason Crow, Jason Smith, Jay Obernolte, Jeanne Shaheen, Jeff Duncan, Jeff Jackson, Jeff Merkley, Jeff Van Drew, Jen Kiggans, Jennifer McClellan, Jennifer Wexton, Jerry Carl, Jerry Moran, Jerry Nadler, Jill Tokuda, Jim Baird, Jim Banks, Jim Clyburn, Jim Costa, Jim Himes, Jim Jordan, Jim McGovern, Jimmy Gomez, Jimmy Panetta, Joaquin Castro, Jodey Arrington, Joe Courtney, Joe Manchin, Joe Neguse, Joe Wilson, John B. Larson, John Barrasso, John Boozman, John Carter, John Cornyn, John Curtis, John Duarte, John Fetterman, John Garamendi, John Hoeven, John James, John Joyce, John Kennedy, John Moolenaar, John Rose, John Rutherford, John Sarbanes, John Thune, John W. Hickenlooper, Jon Ossoff, Jon Tester, Jonathan Jackson, Joni Ernst, Joseph Morelle, Josh Brecheen, Josh Gottheimer, Josh Harder, Josh Hawley, Joyce Beatty, Juan Ciscomani, Juan Vargas, Judy Chu, Julia Brownley, Julia Letlow, Kat Cammack, Katherine Clark, Kathy Castor, Kathy Manning, Katie Boyd Britt, Katie Porter, Kay Granger, Keith Self, Kelly Armstrong, Ken Calvert, Kevin Cramer, Kevin Hern, Kevin Kiley, Kevin Mullin, Kim Schrier, Kirsten E. Gillibrand, Kyrsten Sinema, Lance Gooden, Laphonza R. Butler, Larry Bucshon, Laurel Lee, Lauren Boebert, Lauren Underwood, Linda Sánchez, Lindsey Graham, Lisa Blunt Rochester, Lisa McClain, Lisa Murkowski, Lizzie Fletcher, Lloyd Doggett, Lloyd Smucker, Lois Frankel, Lori Chavez-DeRemer, Lori Trahan, Lou Correa, Lucy McBath, Madeleine Dean, Marc Molinaro, Marc Veasey, Marco Rubio, Marcy Kaptur, Margaret Wood Hassan, Maria Cantwell, María Elvira Salazar, Mariannette Miller-Meeks, Marie Gluesenkamp Perez, Marilyn Strickland, Mar-

jorie Taylor Greene, Mark Alford, Mark Amodei, Mark DeSaulnier, Mark E. Green, Mark Kelly, Mark Pocan, Mark R. Warner, Mark Takano, Markwayne Mullin, Marsha Blackburn, Martin Heinrich, Mary Gay Scanlon, Mary Miller, Mary Peltola, Matt Cartwright, Matt Gaetz, Matt Rosendale, Max Miller, Maxine Waters, Maxwell Frost, Mazie Hirono, Melanie Stansbury, Michael C. Burgess, Michael Cloud, Michael F. Bennet, Michael Guest, Michael McCaul, Michael Rulli, Michael Waltz, Michelle Steel, Mike Bost, Mike Braun, Mike Carey, Mike Crapo, Mike Ezell, Mike Flood, Mike Garcia, Mike Johnson, Mike Kelly, Mike Lawler, Mike Lee, Mike Levin, Mike Quigley, Mike Rogers, Mike Rounds, Mike Simpson, Mike Thompson, Mike Turner, Mikie Sherrill, Mitch McConnell, Mitt Romney, Monica De La Cruz, Morgan Luttrell, Morgan McGarvey, Nancy Mace, Nanette Barragán, Nathaniel Moran, Neal Dunn, Nick LaLota, Nick Langworthy, Nicole Malliotakis, Nikki Budzinski, Norma Torres, Nydia Velázquez, Pat Ryan, Patrick McHenry, Patty Murray, Paul Gosar, Paul Tonko, Pete Aguilar, Pete Ricketts, Pete Sessions, Peter Welch, Pramila Jayapal, Raja Krishnamoorthi, Ralph Norman, Rand Paul, Randy Feenstra, Randy Weber, Raphael G. Warnock, Rashida Tlaib, Raúl Grijalva, Raul Ruiz, Rich McCormick, Richard Blumenthal, Richard Hudson, Richard J. Durbin, Richard Neal, Rick Allen, Rick Crawford, Rick Larsen, Rick Scott, Ritchie Torres, Ro Khanna, Rob Menendez, Robert Aderholt, Robert Garcia, Robert Menendez, Robert P. Casey, Robin Kelly, Roger F. Wicker, Roger Marshall, Roger Williams, Ron Estes, Ron Johnson, Ron Wyden, Ronny Jackson, Rosa DeLauro, Ruben Gallego, Rudy Yakym, Russ Fulcher, Russell Fry, Ryan Zinke, Salud Carbajal, Sam Graves, Sanford Bishop, Sara Jacobs, Scott DesJarlais, Scott Fitzgerald, Scott Franklin, Scott Perry, Scott Peters, Sean Casten, Seth Magaziner, Seth Moulton, Sharice Davids, Sheila Cherfilus-McCormick, Sheldon Whitehouse, Shelley Moore Capito, Sherrod Brown, Shontel Brown, Shri Thanedar, Steny Hoyer, Stephanie Bice, Stephen F. Lynch, Steve Cohen, Steve Daines, Steve Scalise, Steve Womack, Steven Horsford, Summer Lee, Susan M. Collins, Susan Wild, Susie Lee, Suzan DelBene, Suzanne Bonamici, Sydney Kamlager-Dove, Sylvia Garcia, Tammy Baldwin, Tammy Duckworth, Ted Budd, Ted Cruz, Ted Lieu, Teresa Leger Fernandez, Terri Sewell, Thom Tillis, Thomas Kean Jr., Thomas Massie, Thomas R. Carper, Tim Burchett, Tim Kaine, Tim Kennedy, Tim Scott, Tim Walberg, Tina Smith, Todd Young, Tom Cole, Tom Cotton, Tom Emmer, Tom McClintock, Tom Suozzi, Tom Tiffany, Tommy Tuberville, Tony Cárdenas, Tracey Mann, Trent Kelly, Troy Balderson, Troy Carter, Troy Nehls, Val Hoyle, Valerie Foushee, Vern Buchanan, Veronica Escobar, Vicente Gonzalez, Victoria Spartz, Vince Fong, Virginia Foxx, Warren Davidson, Wesley Hunt, Wiley Nickel, William Timmons, Yadira Caraveo, Young Kim, Yvette Clarke, Zach Nunn, Zoe Lofgren

We use the following surrounding keywords as secondary keywords that allow us to ensure that political issues are discussed in a political context:

- Representative, conservative, debate, voters, campaign, liberal, vote, elections, independent, rally, progressive, rights, voting, contribute, debates, election, campaign, republican, fundraising, nomination, candidate, polls, electoral, ballot, reelection, democrat, democracy, suppression, elect, electorate, polling, republicans, nominee, lobbying, constituent, turnout, endorsement, farright, recount, political, recount, absentee, redistricting, populist, progressives, gerrymandering, farleft, president, presidential

We use the following keywords to compare celebrities to politicians. They represent the 100 celebrities with the highest exposure in our data:

- Taylor Swift, Drake, Sabrina Carpenter, Eminem, Billie Eilish, Tyler The Creator, The Weeknd, Lana Del Rey, Travis Scott, Kendrick Lamar, Bruno Mars, Cristiano Ronaldo, Michael Jackson, Justin Bieber, Coldplay, Rihanna, Charli XCX, Kanye West, Linkin Park, Adele, SZA, Katy Perry, Olivia Rodrigo, Doja Cat, Frank Ocean, Shakira, Metro Boomin, Nirvana, Beyoncé, Selena Gomez, Nicki Minaj, Hozier, Imagine Dragons, J Cole, Megan Thee Stallion, Snoop Dogg, Post Malone, Gracie Abrams, Stray Kids, Ed Sheeran, Neymar, Dua Lipa, Chris Brown, Ice Spice, Deftones, Cardi B, Harry Styles, Jennifer Lopez, Gunna, MrBeast, Tate McRae, Kylian Mbappé, Zach Bryan, Fleetwood Mac, Markiplier, Karol G, LeBron James, Morgan Wallen, Benson Boone, Miley Cyrus, Zendaya, Noah Kahan, Logan Paul, Jason Derulo, Jimin, JoJo Siwa, GloRilla, OneRepublic, Virat Kohli, The Kid LAROI, Demi Lovato, Tommy Richman, Jelly Roll, Gordon Ramsay, Alia Bhatt, SIMI, Shraddha Kapoor, Kylie Jenner, Addison Rae, Marques Brownlee, Teddy Swims, Priyanka Chopra, PewDiePie, Bryson Tiller, Mark Rober, Jack Harlow, Cody Johnson, Brandon Lake, Kevin Hart, Jacksepticeye, Jeff Nippard, Dude Perfect, Lionel Messi, Ellen Degeneres, Chris Stapleton, Bailey Zimmerman, JaidenAnimations, El Chapo, James Charles

S9.2 Classification of Applications into Categories

Applications are classified into the main app categories as follows:

- **News Apps:** CNN, BBC News, CBC News, Washington Post, Fox News, NPR, NBC NEWS, NYTimes, Google News, ABC News, AP News, Samsung News, WSJ, BBC, KSL, Deseret News, NewsBreak, Ground News, Guardian, TVA Nouvelles, Público, La Presse, HuffPost, DailyWire, Epoch Times, The Atlantic, inshorts, Denver Post, CNBC, Yahoo News, SmartNews
- **Communication Apps:** Telegram, WhatsApp, Messages, Messenger, Message+, Email, Gmail, Yahoo Mail, Outlook, WhatsApp Business, Signal, GroupMe, Slack, Discord, Kick, TextNow, Mensajes, WeChat, Viber

- **Social Apps:** Facebook, Instagram, TikTok, X, Threads, Truth Social, Snapchat, LinkedIn, Tumblr, Reddit, Pinterest, Instagram Lite, Facebook Lite, TikTok Lite, Bluesky
- **Browser Apps:** Chrome, Safari, Opera, Brave, Firefox, Edge, Kiwi Browser, Free Ad-blocker Browser, XBrowser, Samsung Internet, Opera GX, Firefox Nightly, Opera Mini, Silk Browser, Firefox Focus, Chrome Beta, Vivaldi
- **Music & Video Apps:** YouTube, Spotify, YouTube Music, JioSaavn, SoundCloud, Amazon Music, Audify Music Player, Yandex Music, Audible, Spotify Lite, Spotify X, Spotify for Artists, Spotify for Creators, YouTube Music, YouTube Premium, YouTube Pro, YouTube ReVanced, YouTube ReVanced Extended, YouTube Vanced, YouTube.com, Rumble, Podcast Addict, CleanTube, NewPipe, Netflix, Disney+, Hulu, Prime Video
- **Search Apps:** Google, DuckDuckGo, Ecosia, Bing
- **AI Apps:** ChatGPT, Perplexity, Gemini, Claude, Copilot, Character.AI