

FOR RELEASE AUGUST 27, 2026

No Easy Fix for Bogus Respondents in Online Opt-In Polls

In a test of three screening methods, voter file matching increased error by removing valid cases

BY *Andrew Mercer*

FOR MEDIA OR OTHER INQUIRIES:

Andrew Mercer, Principal Methodologist
Nida Asheer, Senior Communications Manager

202.419.4372
www.pewresearch.org

RECOMMENDED CITATION

Pew Research Center, August 2026, "No Easy Fix for Bogus Respondents in Online Opt-In Polls"

About Pew Research Center

Pew Research Center is a nonpartisan, nonadvocacy fact tank that informs the public about the issues, attitudes and trends shaping the world. It does not take policy positions. The Center conducts public opinion polling, demographic research, computational social science research and other data-driven research. It studies politics and policy; news habits and media; the internet and technology; religion; race and ethnicity; international affairs; social, demographic and economic trends; science; research methodology and data science; and immigration and migration. Pew Research Center is a subsidiary of The Pew Charitable Trusts, its primary funder.

© Pew Research Center 2026

About this research

This study was designed to measure the impact of bogus respondents on opt-in surveys. It also compares three methods for identifying and removing bogus respondents: 1) the use of trap questions, sometimes called “attention checks,” 2) the Sentry prescreening system proprietary to CloudResearch, and 3) matching respondents to a national voter file.

Why did we do this?

Pew Research Center does high-quality research to help the public, the media and decision-makers understand important topics. Our methodological research investigates the current challenges facing the polling industry, including past work on the impact of [bogus respondents in opt-in surveys](#).

Learn more [about Pew Research Center](#) and our [methodological research](#).

How did we do this?

We fielded a large online opt-in survey Nov. 14-19, 2024, among 11,114 U.S. adult respondents. Respondents who agreed to provide their name and contact information were matched to a registered voter file after the survey concluded.

We evaluated how each approach for removing bogus cases performed on three data quality metrics: “yea-saying,” or agreeing regardless of what is asked, 2) quality of open-end text responses, and 3) response order effects. We also examined how screening methods substantively affected estimates of voter turnout and vote choice in the 2024 presidential election.

Here are the [questions used for this report](#) the survey [methodology](#).

No Easy Fix for Bogus Respondents in Online Opt-In Polls

In a test of three screening methods, voter file matching increased error by removing valid cases

One of the most urgent problems in online opt-in polling is [bogus \(or fraudulent\) respondents](#). These are survey-takers who make no effort to answer questions truthfully and instead are just looking to finish surveys quickly and collect rewards.

To combat this threat, researchers have developed various ways to identify and purge bogus cases from survey samples. Approaches include 1) trap questions, sometimes called “attention checks,” that genuine respondents should always answer correctly, 2) automated prescreening services, and 3) matching respondents to a registered voter file.

But how well do they work? A new Pew Research Center study finds that:

- Overall, purging bogus cases lowers error on most metrics, but a surefire solution remains elusive.
- Matching an opt-in sample to voter files slightly *increased* error by removing mostly good respondents (e.g., those who simply declined to give their name or address).
- Trap questions and an automated prescreening service performed similarly, improving data quality somewhat.

Matching online opt-in respondents to voter file did not reduce error

Estimated error on 3 data quality measures after using each method for removing bogus respondents

Error measure	No screening	Trap questions	Automated prescreening	Voter file matching
Yea-saying (acquiescence bias)	7%	1% ▼ 6 PTS LESS	2% ▼ 5 PTS LESS	9% ▲ 2 PTS MORE
Problematic open-ends	12%	7% ▼ 5 PTS LESS	5% ▼ 7 PTS LESS	13% ▲ 1 PT MORE
Order effect	+7 pct. pts.	+3 ▼ 4 PTS LESS	+4 ▼ 3 PTS LESS	+9 ▲ 2 PTS MORE

Note: Acquiescence bias was computed as the share of respondents answering “Yes” to at least 10 of 15 questions with yes/no answer choices. Order effect was computed as an average across 13 questions.

Source: Online opt-in survey of U.S. adults conducted Nov. 14-19, 2024.

“No Easy Fix for Bogus Respondents in Online Opt-In Polls”

PEW RESEARCH CENTER

- All three approaches modestly increased the overestimation of Democratic support in the 2024 election. This appears to be due not to a systematic partisan bias but to bogus respondents' tendency to say they voted for the winning candidate – in this case, Donald Trump, the Republican.

Matching opt-in samples to voter files may serve other purposes, such as providing data on respondents' voting history. This study speaks only to whether this is an effective tool for purging bogus cases.

Do all polls have bogus respondents?

No. Bogus respondents are primarily a threat to online opt-in polls, which are recruited through methods like online advertising, self-enrollment and email lists. Polls recruited offline using random sampling (e.g., Pew Research Center's [American Trends Panel](#)) are generally immune to this threat, though they face [other challenges](#).

Related: [*Do AI and bogus respondents threaten polling's future?*](#)

Methods for removing bogus respondents

Studies conducted by [Center researchers](#) and [others](#) have found that standard data quality checks such as looking for “speeders” who complete surveys too quickly or “straightliners” who consistently select the same answer choice (e.g., always pick the first response option) fail to detect most bogus respondents. The same goes for many trap questions or attention checks, which not only fail to detect bogus respondents but can also confuse legitimate respondents, [leading to false positives](#).

Faced with these challenges, pollsters have developed a variety of new approaches in hopes of better identifying bogus respondents.

In this study, we evaluate three such approaches:

- Trap questions about unlikely behaviors
- Automated prescreening by a leading fraud-detection company
- Matching respondents to a commercial voter file

Pros and cons of 3 approaches for removing bogus respondents from opt-in samples

	Trap questions	Automated prescreening	Voter file matching
Pros	• No major added cost	• Researcher does not have to implement • Prescreening companies may stay more current on best practices than researchers	• Requiring a valid voter name and address should be a major barrier for bogus respondents
Cons	• Bad actors are aware of trap questions and, in some cases, know how to circumvent them • Trap questions can degrade the survey experience for good respondents	• Added cost from using prescreening service • Prescreening can degrade the survey experience for good respondents	• Added cost from using voter file • Screens out good respondents who decline matching on privacy grounds • Screens out good respondents who do not happen to be registered voters • This Center study finds voter file matching tends to increase error by disproportionately removing good cases

“No Easy Fix for Bogus Respondents in Online Opt-In Polls”

PEW RESEARCH CENTER

Trap questions

A difficulty in identifying bogus respondents lies in the fact that, for most questions, we have no way of knowing if a respondent has answered truthfully. To get around this problem, we asked respondents if they had ever engaged in four activities where we can be virtually certain that the true answer is “No.” Under this screening procedure, respondents were coded as bogus if they answered “Yes” to one or more of these questions.

The questions were designed so that diligent respondents would not be confused about how to answer, while bogus respondents trying to appear eligible for surveys targeting specific groups or answering randomly would be more likely to answer “Yes.” Specifically:

- We asked respondents if they used any of six different social media platforms, one of which was a made-up platform called Fizzypress. We also asked if they had received any of six different government benefits in 2023, including payments from the [United States Railroad Administration](#) (USRA) – a real government agency, but one that ceased to exist in 1920. Both of these were asked early on in the survey and were intended to resemble screening questions looking for users of specific social media platforms or recipients of certain kinds of benefits. Of the full sample, 6% said they used Fizzypress and 9% claimed to have received USRA payments.
- Later in the survey, we asked respondents if they had ever visited the International Space Station; 10% said they had. Finally, we asked if they had ever served on a [Polar-class icebreaker ship](#), only two of which were ever built and only one of which remains in active use by the U.S. Coast Guard; 8% answered in the affirmative.

A total of 1,963 cases (18%) answered “Yes” to one or more of these trap questions and were flagged as bogus.

Automated prescreening

CloudResearch’s proprietary [Sentry](#) prescreening system is designed to automatically identify problematic respondents before they begin a survey. When a potential respondent first clicks the survey invitation link, they are routed to the Sentry platform and asked a short series of questions designed to elicit problematic survey-taking behaviors like yea-saying and inattentive responding. Open-end answers are checked to confirm that they are responsive to the question asked and not pasted from another source via an automated process.

The system also performs passive checks using metadata about the respondent's device, location and browser for other signs that they are misrepresenting themselves through technical means.

Respondents who pass these checks are then routed to take the survey. Typically, those who fail are not forwarded to the main survey, but for this study, no respondents were terminated for failing the prescreening checks. We used metadata about which cases failed and why to simulate what would have happened to the survey results if they were excluded at the outset.

Under this procedure, 5,369 cases – nearly half of all completes – failed at least one prescreening check, including 1,879 (35% of failed cases) that failed multiple.

- 70% of failed cases had a problematic open-end, making this the most frequently failed check by far. This was followed by yea-saying (47%), inattentiveness (12%) and duplicate IP addresses (10%).
- 7% of these cases failed checks for fraudulent behavior, and only 1% failed other passive technical checks.

Voter file matching

The third approach we tested involves asking respondents to provide their name and address and then looking for a matching record in a [commercial voter file](#), a national database of nearly all registered voters in the United States. If a matching record can be found, the case is considered valid. If no matching record is found, the case is thrown out.

This method assumes that respondents who can be matched are most likely being honest about their identity, and that someone willing to provide detailed contact information that can be validated against official voting records will likely be diligent about answering other survey questions.

This approach has primarily been used by pollsters who [work for political campaigns](#). For cases that are successfully matched, voter files can provide a great deal of information beyond what was asked in the survey, such as a respondent's registration status and voting history.

But while many voter file vendors attempt to include the unregistered population, [Center research](#) has found that a sizable share of this group is not covered by these databases. This can lead to throwing away data for many otherwise-valid respondents who simply aren't registered to vote. This is not a large drawback for political pollsters, who typically focus on surveying registered

voters. But if certain kinds of people who are more likely to be unregistered are underrepresented in a sample or missing altogether, it may lead to biased results if applied to a general population survey.

In this study, we asked respondents if they were willing to provide their name and home address so that they could be matched to the voter file. Out of all respondents, 49% agreed to provide their contact information; of these, 73% were successfully matched to the [TargetSmart](#) voter file based on name, address, age and sex. Altogether, 3,977 cases (36% of the full sample) were successfully matched while 7,137 (64%) were screened out under this procedure.

Distinct demographic patterns in flagged cases

Although the overall number of cases screened out by each method varies dramatically, respondents claiming to have certain demographic characteristics were consistently more likely to be flagged. Specifically, trap questions and prescreening were especially likely to flag those claiming to be ages 18 to 29 or Hispanic – the two groups found to be most affected by bogus responding in previous Center studies.

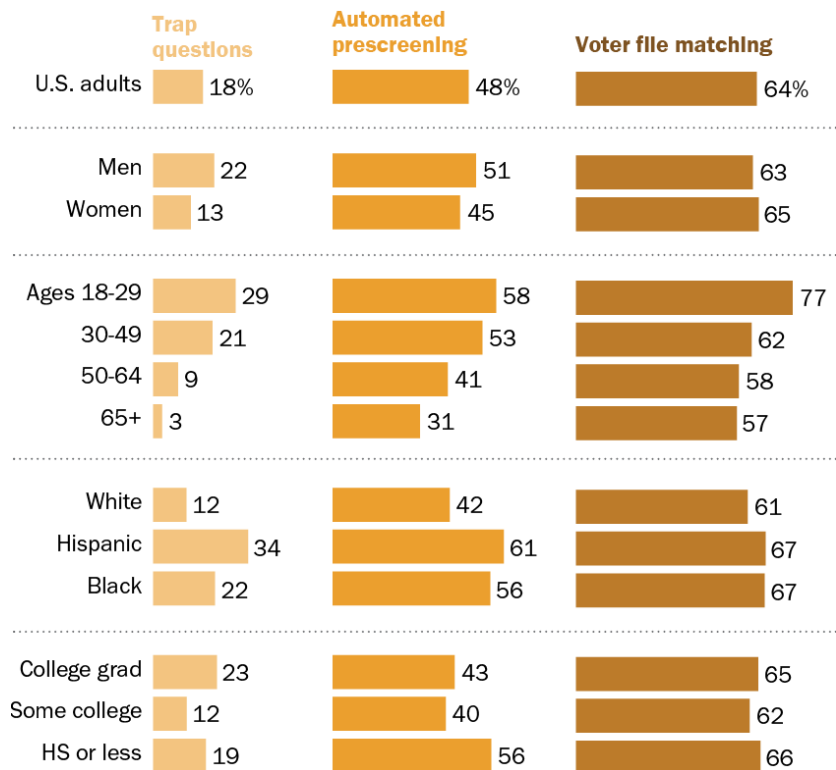
This should not be taken to mean that members of these demographic groups tend to be poor survey-takers. Rather, insincere respondents tend to *claim* membership in these groups, likely falsely in many cases.

The profile of cases purged using voter file matching was different. This reflects the fact that voter file matching tends

to purge cases for benign reasons (e.g., a respondent's privacy concerns or their not being a registered voter) not harmful ones. Hispanic cases were no more likely to be screened out than non-Hispanic Black cases, and both were only 6 percentage points more likely to be screened out than non-Hispanic White cases. In contrast, Hispanic cases were more likely to be flagged than non-Hispanic White cases by trap questions (+22 points) and prescreening (+19).

Across demographic groups, prescreening and voter file matching flag far more cases than trap questions as potentially bogus

% flagged as bogus by each screening method



Note: White and Black adults include those who report being one race and are not Hispanic; Hispanics are of any race.

Source: Online opt-in survey of U.S. adults conducted Nov. 14-19, 2024.

"No Easy Fix for Bogus Respondents in Online Opt-In Polls"

PEW RESEARCH CENTER

Likewise, trap questions and prescreening were both more likely to flag men than women, by margins of 9 and 6 points, respectively. Men and women were screened out by voter file matching at roughly the same rate.

The three methods differed on education:

- With trap questions, college graduates (+11) and respondents with a high school education or less (+7) were more likely to be flagged than those with some college education.
- With prescreening, high school or less cases (+13) were more likely to be flagged than college graduates, who had the next-highest flag rate.
- Voter file matching showed much less differentiation: Rates across education groups did not vary by more than 4 points.

Effect of screening on measures of data quality

What impact do these three screening methods have on data quality? There is no direct way to know what proportion of bogus respondents are successfully identified by each method, nor can we be sure how many valid cases are misidentified as fraudulent. Instead, we can compare various measures of data quality calculated before and after each screening method has been applied and the flagged cases removed. To ensure that results are comparable across screening methods, separate sets of survey weights were created for the unscreened sample and for the set of cases remaining after each method was applied. For details, refer to the [methodology](#).

In this study, we focused on three measures of data quality: 1) the frequency answering “Yes” to yes/no questions, sometimes called “yea-saying,” 2) the quality of text responses to open-ended questions, and 3) the severity of response order effects when the order of answer choices is randomized.

Yea-saying

In a [2023 Pew Research Center study](#), one indicator of bogus responding was a tendency to answer yes/no questions in the affirmative, regardless of the question. In that study, 8% of online opt-in respondents answered “Yes” to at least 10 of 16 yes/no questions asked, with especially high shares among 18- to 29-year-olds (15%) and Hispanic respondents (19%). In contrast, the corresponding shares among respondents from probability-based panels fell between 1% and 2%. Many of these questions measured rare or uncommon characteristics, and we would expect that virtually no one answering truthfully would say “Yes” to 10 or more.

Our latest survey exhibits largely the same pattern. Prior to any screening, 7% of all adults, 10% of 18- to 29-year-olds and 13% of Hispanic respondents answered “Yes” to at least 10 of 15 yes/no questions (excluding a question measuring Hispanic ethnicity and questions used in any screening procedures).

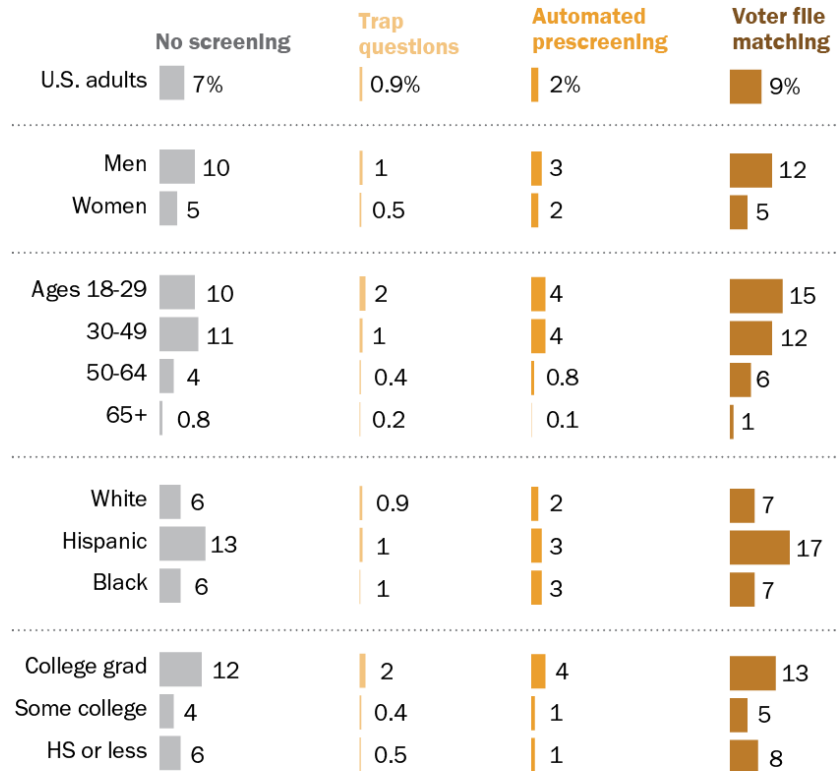
After screening with trap questions, this behavior is largely eliminated from the remaining respondents: 1% of adults and no more than 2% of any demographic subgroup answered 10 or more yes/no questions in the affirmative. The reduction is almost as large for prescreening, which reduced the shares to 2% for all adults, 4% for young adults and 3% for Hispanic adults.

In contrast, matching to the voter file appears to have made the problem worse. The share among all adults increased slightly to 12%, and the shares among young and Hispanic adults rose to 25% and 18%, respectively.

This surprising result appears to be due to two factors. The first is that respondents who declined to provide their contact information – nearly half of the sample – appear much less likely to be bogus than those who agreed to provide this information. Among those who declined, only 3% answered “Yes” to 10 or more yes/no questions, compared with 12% among those who agreed.

Trap questions nearly eliminate yea-saying while voter file matching makes it worse

% who answered “Yes” to 10 or more yes/no questions after applying each screening method



Note: White and Black adults include those who report being one race and are not Hispanic; Hispanics are of any race.

Source: Online opt-in survey of U.S. adults conducted Nov. 14-19, 2024.

“No Easy Fix for Bogus Respondents in Online Opt-In Polls”

PEW RESEARCH CENTER

The second is that a nontrivial share of bogus respondents were able to provide contact information that was accurate and detailed enough to result in a successful voter file match, though the information provided may not be their own. To the extent that some bogus respondents were successfully removed, it was not enough to offset the loss of valid respondents who declined to share contact information.

Open-end response quality

Examining text answers to open-ended questions has long been considered a reliable way to identify low-quality respondents. While attentive respondents give answers that are relevant to the question asked, bogus respondents often give [nonsensical or gibberish answers](#).

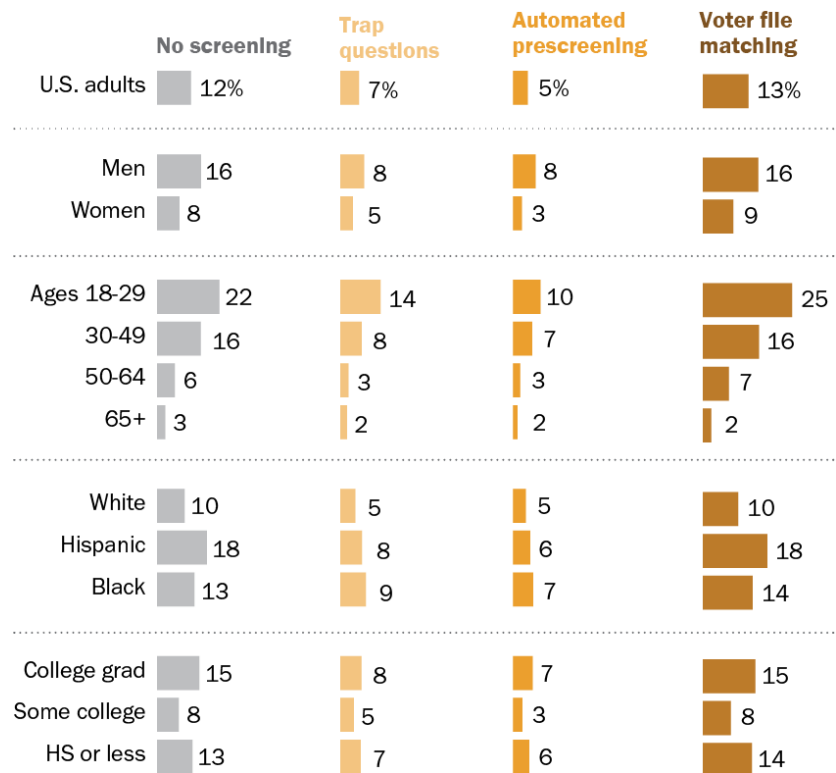
This survey included an open-ended question that asked, “What is one thing that you would like politicians in Washington, D.C. to know about your own situation when they are writing laws and setting policy? Please share as much detail as you can.”

Each respondent’s answer was reviewed and assigned to one of the following categories: 1) relevant, 2) nonresponse, 3) generic positive rating, 4) other non sequitur, 5) gibberish, or 6) probable AI (refer to the [methodology for details](#)). Because respondents had been told they could skip any question they did not wish to answer, both relevant and nonresponse answers were considered unproblematic, while the remaining categories were deemed problematic.

Prior to any screening, 88% of answers were classified as unproblematic, including 69% that were relevant and 19% that were nonresponse. The most common kind of problematic answers were other non sequiturs (7%), followed by generic positive

Voter file matching had minimal effect on open-end response quality

% who gave problematic open-end answers after applying each screening method



Note: White and Black adults include those who report being one race and are not Hispanic; Hispanics are of any race.

Source: Online opt-in survey of U.S. adults conducted Nov. 14-19, 2024.

“No Easy Fix for Bogus Respondents in Online Opt-In Polls”

PEW RESEARCH CENTER

ratings (2%), gibberish (2%) and probable AI (1%). It is worth noting that the probable AI responses detected were written in a distinct style that was relatively easy to spot. There may be other, more subtle, AI responses that were not detected.

Among all adults, trap questions and prescreening performed similarly, bringing the overall share of problematic open-ends down to 7% and 5%, respectively. Prescreening also resulted in a higher share of relevant answers than trap questions (80% vs. 74%) and a lower share of nonresponse (15% vs. 19%). This may be because the prescreening process itself includes an open-ended question. Matching to voter files, on the other hand, had virtually no effect on the proportion of problematic answers, though the proportion of relevant answers increased to 73%.

Response order effects

If a respondent is being attentive and answering questions accurately, the answer options they select shouldn't be affected by the order in which they are presented. But inattentive respondents in online surveys tend to select answer choices that appear toward the top of the list, otherwise known as a "primacy effect."

Our survey included a question asking if undocumented immigrants should be allowed to remain in the country, after which a random half of respondents were shown answer choices in the following order:

- *"They should not be allowed to stay in the country legally."*
- *"There should be a way for them to stay in the country legally, if certain requirements are met."*

For the other half of respondents, this order was reversed. If all respondents were answering diligently, the results from each half of the sample would be largely the same.

Instead, we saw a substantial primacy effect on this question. Without any screening, 43% of adults endorse the view that undocumented immigrants should not be allowed to stay legally when that option is presented first. The share drops to 34% when the order is reversed, a primacy effect of +9 percentage points.

The primacy effects were even larger for those subgroups most affected by bogus

responding, at +13 points for men, +15 points for 18- to 29-year-olds and +17 points for Hispanic adults.

Trap questions and prescreening reduced primacy effects on immigration question; voter file matching made them worse

% who say undocumented immigrants should not be allowed to stay in the country legally when that response option is shown first vs. second

	No screening			Primacy effect after screening (1st-2nd)		
	Shown 1st	Shown 2nd	Primacy effect (1st-2nd)	Trap questions	Pre- screening	Voter file matching
All U.S. adults	43%	34%	+9	+4	+5	+14
Men	49	36	+13	+6	+8	+18
Ages 18-29	42	27	+15	+7	+10	+29
Hispanic	44	27	+17	+9	+7	+25

Source: Online opt-in survey of U.S. adults conducted Nov. 14-19, 2024.

"No Easy Fix for Bogus Respondents in Online Opt-In Polls"

PEW RESEARCH CENTER

- Screening with trap questions reduced the primacy effect by roughly half, to +4 points for all adults, +6 for men, +7 for young adults and +9 for Hispanic adults.
- Prescreening performed similarly, reducing the primacy effect to between +5 and +10 points.
- Voter file matching, on the other hand, resulted in *even larger* primacy effects than when there was no screening at all, with magnitudes increasing to +14 for all adults, +18 for men, +29 for young adults and +25 for Hispanic adults.

As with the frequency of giving "Yes" answers, voter file matching's negative impact on data quality appears to be due to valid respondents choosing not to provide contact information for matching, with those cases showing a primacy effect of only +2 points.

Including the question about the status of undocumented immigrants, there were a total of 13 questions with randomized response options that were asked of all respondents. Nearly all of these showed the same pattern. Without any screening, the average primacy effect on these questions was +7 percentage points for all adults, and as high as +11 and +12 points for young adults and Hispanic adults, respectively. Trap questions reduced the average primacy effect by 4 points for all adults and by 7 points for young and Hispanic adults, while prescreening did almost as well.

Voter file matching made primacy effects worse

Average primacy effect across 13 questions with randomized response options

	No screening	Trap questions	Pre- screening	Voter file matching
All U.S. adults	+7	+3	+4	+9
Men	+9	+4	+5	+11
Women	+5	+3	+3	+7
Ages 18-29	+11	+4	+6	+15
30-49	+10	+5	+5	+11
50-64	+4	+2	+3	+5
65+	+3	+2	+2	+2
White	+6	+3	+4	+6
Hispanic	+12	+5	+6	+16
Black	+8	+4	+5	+10
College grad	+7	+3	+4	+11
Some college	+5	+3	+4	+6
HS or less	+8	+3	+4	+9

Note: Primacy effect is defined as the difference between the percent who selected a category when it is presented first minus the percent who selected it when the order was reversed. For ordered scales, the top two and bottom two categories were collapsed. Because individual questions can have multiple primacy effects, the largest primacy effect for each question is used to compute averages. White and Black adults include those who report being one race and are not Hispanic; Hispanics are of any race.

Source: Online opt-in survey of U.S. adults conducted Nov. 14-19, 2024.
"No Easy Fix for Bogus Respondents in Online Opt-In Polls"

PEW RESEARCH CENTER

How does screening affect questions about voter turnout and vote choice?

The prevalence of online opt-in samples in election polling raises the question of how different screening methods – voter file matching in particular – impact estimates related to voter turnout and vote choice. Fielded shortly after the 2024 U.S. presidential election, this survey asked respondents if they voted, and if so, for whom.

Without any screening, an estimated 84% of self-reported registered voters said they voted in the 2024 presidential election. Purging cases with trap questions and prescreening both increased that share by 3 points, to 87%, while voter file matching increased it by 1 point, to 85%. This pattern is more pronounced among young adults and Hispanic adults: For these groups, trap questions and prescreening increased estimated turnout by between 6 and 9 percentage points, while voter file matching produced a 1-point increase among young adults and a 2-point decrease among Hispanic adults.

Although an increase in voter turnout would appear to make these estimates *less* accurate relative to a higher quality estimate of roughly 77%, it is notable that, for this question, the order of response options was not randomized and the response option for having voted was presented last.¹ This kind of pattern is what we would expect to see if bogus respondents, who are more likely to select answer choices toward the top of the list, are being screened out. The fact that the effect is smaller for voter file matching is also consistent with that method's tendency to disproportionately exclude valid respondents.

All three screening methods increased estimated support for Harris

When it comes to presidential vote (for which response option order *was* randomized), the weighted raw sample showed Donald Trump and Kamala Harris tied, each with 48% of the vote – accurate to Harris' true vote share, but underestimating Trump's by 2 points. Trap questions slightly shifted the result from a tie to a 3-point advantage for Harris. The effect was larger for prescreening and voter file matching, which shifted Harris' margin to +6 and +7 points, respectively.

These patterns suggest that Harris supporters are overrepresented among valid respondents. But prior to screening, this bias was largely offset by the presence of bogus respondents, who disproportionately said they supported Trump. When bogus respondents were screened out, the bias in favor of Harris became more apparent.

¹ 77% voter turnout is based on an analysis of the validated voter dataset from Pew Research Center's [2024 postelection survey](#).

Because the answer choices for Trump and Harris were randomized, this result also implies that bogus respondents were not simply choosing the first option or choosing randomly. There appears to be a subset of bogus respondents, specifically those with the most glaring data quality problems, who were much more likely to say that they voted for Trump, regardless of whether that option was presented first or second. For example, in the unscreened sample, cases coded as having a problematic open-end supported Trump over Harris by a margin of 29 percentage points, while cases that gave 10 or more “Yes” answers favored Trump by 30 points.

Cases with one or both of these data quality problems made up anywhere from 11% of screenouts from voter file matching up to 57% of screenouts from trap questions.

This should not be taken to mean that bogus respondents will always be systematically biased in favor of Republican candidates. In a [previous Center benchmarking study](#) where respondents were asked about the 2020 presidential election, opt-in respondents who gave 10 or more “Yes” answers claimed to have voted for Joe Biden, the Democrat, over Trump by margins ranging from 34 to 51 points across three different opt-in samples. It seems plausible that, when asked about past elections, many of these respondents are simply choosing the candidate who won.

These effects are even larger for subgroups in which bogus respondents are most common. Among 18- to 29-year-old voters, the survey initially showed Trump ahead by 2 points prior to screening. All three screening approaches shifted the margin in favor of Harris, putting her ahead by 3 points with trap questions, 6 points with prescreening and 1 point with voter file matching. Among Hispanic voters, trap questions slightly decreased Harris’ lead from 9 points to 8. Prescreening and voter file matching had larger effects, increasing Harris’ lead among Hispanic voters to 16 and 19 points, respectively.

Appendix: What are bogus respondents?

Survey researchers have long been concerned about respondents who answer questions inattentively. For traditional probability-based surveys, such respondents typically make up only a small fraction of those sampled. The process of selecting respondents at random from a list containing the entire population of interest means there is no way for potential bad actors to self-select into probability-based samples.

For online opt-in samples, individuals can sign up to take as many surveys as they want, collecting a small reward for each one. This creates an incentive for people or organizations to take as many surveys as possible in order to maximize their reward. As a result, online opt-in samples have been flooded with “[bogus respondents](#)” who make no effort to answer questions truthfully and instead try to complete surveys as quickly as possible without being disqualified.

The presence of bogus respondents in opt-in samples can lead to large survey errors, particularly on questions about [rare attitudes and behaviors](#). Bogus respondents also disproportionately claim to be either young or Hispanic, resulting in dramatic survey errors for these groups that were, on average, nearly twice as large as the error for the general population in a [2023 Pew Research benchmarking study](#).

This could be because young and Hispanic respondents are highly sought-after by researchers, and by claiming to be members of these groups, bogus respondents increase their chances of qualifying for surveys. It could also be due to question format. For example, Hispanic ethnicity is frequently measured using a yes/no question, and bogus respondents [have been found](#) to answer such questions in the affirmative – regardless of their topic.

Acknowledgments

This report is a collaborative effort based on the input and analysis of the following individuals.

Research team

Andrew Mercer, *Principal Methodologist*

Arnold Lau, *Research Methodologist*

Courtney Kennedy, *former Vice President of Methods and Innovation*

Communications and editorial

Nida Asheer, *Senior Communications Manager*

DeVonte Smith, *Communications Associate*

Anna Jackson, *Editorial Specialist*

Graphic design and web publishing

Bill Webster, *Senior Information Graphics Designer*

Sara Atske, *Digital Producer*

Ashley Amaya and Scott Keeter also provided helpful input on this study.

Methodology

The data in this report comes from an online opt-in survey of 11,114 U.S. adults conducted Nov. 14-19, 2024. The sample was provided by CloudResearch and originally sourced from Cint’s Lucid Marketplace. Adults ages 18 and older located in the United States were eligible to participate. The survey did not use any demographic quotas.

The full set of respondents was then screened using three different methods designed to identify bogus or other problematic respondents: 1) trap questions, 2) automated prescreening using CloudResearch’s proprietary Sentry system, and 3) voter file matching. For each method, we identified which cases passed the checks and which would have been screened out.

Out of the 11,114 respondents overall, a total of 9,151 passed the trap questions, 5,745 passed the automated prescreening and 3,977 were successfully matched to a voter file.

Survey weights were created for the three groups of respondents that would have resulted if each of the three screening methods had been used to remove problematic respondents. A fourth set of weights was created for the entire sample with no screening applied.

Because probabilities of selection are unknown for opt-in samples, respondents were all treated as if their probabilities of selection were all equal and given an initial weight of 1. For each set of respondents, this initial weight was then calibrated to align with the population benchmarks identified in the accompanying table and trimmed at the 1st and 99th percentiles.

Weighting dimensions

Variable	Benchmark source
Age (detailed)	2023 American Community Survey (ACS)
Age x Gender	
Education x Gender	
Education x Age	
Race/Ethnicity x Education	
Race/Ethnicity x Gender	
Race/Ethnicity x Age	
Census region x Metropolitan status	

Note: Estimates from the ACS are based on noninstitutionalized adults.

PEW RESEARCH CENTER

Open-end coding

One measure of data quality examined in this report was whether respondents answers to an open-ended question was meaningful and responsive to the question that was asked. Respondents were asked, *“What is one thing that you would like politicians in Washington, D.C. to know about your own situation when they are writing laws and setting policy? Please share as much detail as you can.”*

Each respondent’s answer was assigned to one the following categories:

- **Relevant:** Answers that were responsive to the question, describing some aspect of the respondent’s life or broadly related to politics.
- **Nonresponse:** Answers indicating there was nothing the respondent wanted politicians to know or that they were unsure. It also included those who declined to answer the question either by stating as much or by skipping the question.
- **Generic positive rating:** Answers that seem like they are giving a postive rating to a product or service, such as “Very good,” “Excellent quality” or “Good service.” It also included uncontextualized positive answers like “Yes” or “Okay.”
- **Other non sequitur:** Answers that do not make sense in the context of the question, such as “Open a Checking Account Online for Free,” “Pizza” or “Hey I will call you in about an hour.”
- **Gibberish:** Answers that look like someone was mashing keys on a keyboard, such as “kljasdfjkd” or “V6f66ff6 6v6x5xuv c5d6f6g6f5 v6f6 cc 6v.”
- **Probable AI:** Answers that, rather than expressing genuine respondent opinion, appeared to have been generated by AI systems like ChatGPT or Gemini. Sometimes this was explicit, such as when the answer includes phrases like “As an AI ...” or “If I had a personal situation to share ...”. Other times, these were answers exhibiting an overly formal style that appeared highly unlikely to have been written by a human.²

This coding scheme was adapted from a [codebook developed](#) for a previous Pew Research Center study of bogus respondents. To validate the reliability of the codebook, a random sample of 300 answers were independently coded by two human reviewers. The sample was stratified on the

² The probable AI category was added to the codebook after the human-coded benchmarks had already been created. These answers were initially categorized as relevant and were treated as such when evaluating the LLM’s accuracy in coding open-ends.

number of incorrect answers each respondent gave to three trap questions to ensure that there was a sufficient number of answers across a range of data quality levels. The double coded answers yielded a Krippendorff's alpha of 0.85 – a level generally considered to indicate a high degree of reliability.

Another 900 similarly sampled answers were then coded and combined with the initial validation sample to create a benchmarking set of 1,200 human-coded answers considered to be coded correctly and serve as a ground truth. We developed a prompt for GPT-4o, providing it with context about the task and criteria for how to assign codes. We then used it to code the benchmarking dataset, evaluated the answers where the human- and large-language-model (LLM)-coded answers differed, and refined the prompt to correct any systematic problems. This process was repeated several times until the overall level of agreement was satisfactory.

The final prompt yielded an agreement rate of 89% and an F1 score of 0.86, indicating generally high levels of classification accuracy for the variable as a whole. Individual categories had agreement rates between 91% and 98% and F1 scores between 0.92 and 0.99.

The final prompt was then applied to the remaining answers to provide the codes used in this analysis. Refer to the [final prompt](#).

Quality metrics for LLM-coded open-ends

Agreement rate and F1 score for LLM-coded benchmark dataset

Category	% agree	F1
All categories	89%	0.86
Relevant/ Probable AI	93	0.92
Nonresponse	98	0.99
Generic positive rating	98	0.99
Other non sequitur	91	0.94
Gibberish	98	0.99

Note: The agreement rate is the percentage of cases in which the LLM's code were the same. The F1 score is the harmonic mean of the model's precision and recall. The probable AI category was added to the codebook after the human-coded benchmarks were created. These answers were recoded as relevant for purposes of computing accuracy metrics.

PEW RESEARCH CENTER

2024 POSTELECTION OPT-IN DATA QUALITY SURVEY**Nov. 14-19, 2024****SURVEY QUESTIONNAIRE****TXT: INTRO****ASK ALL:**

Welcome! Thank you for participating in our survey. It should take about 10 minutes for most people to complete. Here are some helpful hints:

- The information you provide will be used for research purposes only.
- You are not required to answer any question you do not wish to answer. You can click on the Next button to skip a question you would not like to answer.
- Please do not use your browser's back button to go back to previous questions. Instead, use the navigation buttons on each web page to move through the survey.

Click or tap on the Next button to continue.

QUE: YOBTMOD**ASK ALL:**

[PN: PROGRAM AS DROP-DOWN MENU, WITH YEARS 2006-1919 IN DESCENDING ORDER]

In what year were you born?

[PN: PROGRAM DROP-DOWN MENU]

QUE: EDUC_ACS**ASK ALL:**

What is the highest degree or level of school that you have completed?

- 1 No schooling completed
 - 2 Nursery school
 - 3 Kindergarten
 - 4 Grade 1 through 11
 - 5 12th Grade – NO DIPLOMA
 - 6 Regular high school diploma
 - 7 GED or alternative credential
 - 8 Some college credit, but less than 1 year of college credit
 - 9 1 or more years of college credit, no degree
 - 10 Associate degree (for example: AA, AS)
 - 11 Bachelor's degree (for example: BA, BS)
 - 12 Master's degree (for example: MA, MS, MEng, MEd, MSW, MBA)
 - 13 Professional degree beyond a bachelor's degree (for example: MD, DDS, DVM, LLB, JD)
 - 14 Doctorate degree (for example: PhD, EdD)
-

QUE: HISP**ASK ALL:**

Are you of Hispanic, Latino, or Spanish origin, such as Mexican, Puerto Rican, or Cuban?

- 1 Yes
- 2 No

QUE: RACEMOD**ASK ALL:**

What is your race or origin?

[Check all that apply]

- 1 White
- 2 Black or African American
- 3 Asian or Asian American
- 4 American Indian or Alaska Native
- 5 Native Hawaiian or other Pacific Islander
- 6 Some other race or origin (please specify): [PN: INSERT SINGLE LINE TEXT BOX]

QUE: GENDER**ASK ALL:**

Do you describe yourself as a man, a woman, or in some other way?

- 1 A man
- 2 A woman
- 3 In some other way

QUE: ZIP**ASK ALL:**

What is your zip code?

[PN: NUMERIC TEXT BOX]

BAT: SMUSE**ASK ALL:**

Please indicate whether or not you ever use the following websites or apps.

[PN: RANDOMIZE ITEMS]

BATTERY ITEMS:

- a. Facebook
- b. YouTube
- c. X (formerly Twitter)
- d. Instagram
- j. BeReal
- z. Fizzypress

RESPONSE OPTIONS

- 1 Yes, use this
- 2 No, don't use this

QUE: JOBLASTYR_CPS**ASK ALL:**

These next few questions are about things that happened *last year in 2023*.

Did you work at a job or business at any time during 2023?

- 1 Yes
 - 2 No
-

BAT: BENEFITS_CPS**ASK ALL:**

At any time during 2023 did you receive any of the following...

[PN: RANDOMIZE ITEMS]

BATTERY ITEMS:

- a. State or Federal unemployment compensation
- b. Workers' Compensation payments or other payments as a result of a job-related injury or illness
- c. Social Security payments from the U.S. Government
- d. Benefits from SNAP (the Supplemental Nutritional Assistance Program) or the Food Stamp program, or use a SNAP or food stamp benefit card
- e. Payments from the United States Railroad Administration
- f. Veterans' (VA) payments

RESPONSE CATEGORIES:

- 1 Yes
 - 2 No
-

Thinking about the recent presidential election...

QUE: VOTED**ASK ALL:**

Which of the following statements best describes you?

- 1 I did not vote in the November 2024 election
 - 2 I planned to vote but wasn't able to
 - 3 I definitely voted in the November 2024 election
-

QUE: VOTEGEN_POST**ASK IF VOTED (VOTED=3):**

In the 2024 presidential election, who did you vote for?

[PN: RANDOMIZE OPTIONS 1 AND 2]

- 1 Donald Trump, the Republican
- 2 Kamala Harris, the Democrat
- 3 Robert F. Kennedy Jr., a third-party candidate
- 4 Chase Oliver, the Libertarian Party candidate
- 5 Jill Stein, the Green Party candidate
- 6 Cornel West, a third-party candidate

- 7 Claudia De la Cruz, the Socialism and Liberation Party candidate
 8 Another candidate
-

QUE: VOTEGEN_POSTNON**ASK IF DID NOT VOTE (VOTED=1-2):**

If you had voted in the presidential election, would you have voted for...

[PN: ROTATE RESPONSE OPTIONS 1 AND 2]

- 1 Donald Trump, the Republican
 2 Kamala Harris, the Democrat
 3 Robert F. Kennedy Jr., a third-party candidate
 4 Chase Oliver, the Libertarian Party candidate
 5 Jill Stein, the Green Party candidate
 6 Cornel West, a third-party candidate
 7 Claudia De la Cruz, the Socialism and Liberation Party candidate
 8 Another candidate
-

QUE: DTFORAGNST_POST**ASK IF VOTED TRUMP (VOTEGEN_POST=1):**

Would you say that your vote for Trump was more a vote...

- 1 For Donald Trump
 2 Against Kamala Harris
-

QUE: KHFORAGNST_POST**ASK IF VOTED HARRIS (VOTEGEN_POST=2):**

Would you say that your vote for Harris was more a vote...

- 1 For Kamala Harris
 2 Against Donald Trump
-

QUE: VOTEDECTIME**ASK IF VOTED (VOTED=3):**

As far as you can remember, when did you make up your mind about who you were going to vote for in the presidential election?

- 1 Last few days before Election Day
 2 The last week before Election Day
 3 In October
 4 In September
 5 Before September
-

QUE: CONG_POST**ASK IF VOTED (VOTED=3):**

In the 2024 election for U.S. House of Representatives, did you vote for...

[PN: RANDOMIZE OPTIONS 1 AND 2 FIRST FOLLOWED BY 3 AND 4 IN ORDER, WITH SPACE BETWEEN 3 AND 4]

- 1 The Republican Party's candidate
- 2 The Democratic Party's candidate
- 3 Another candidate [ANCHOR]
- [SPACE]
- 3 I did not cast a vote for the House of Representatives in the election [ANCHOR]

QUE: FIRST**ASK IF VOTED (VOTED=3):**

Is this the first year you have ever voted, or have you voted in elections before this year?

- 1 First year voting
- 2 Have voted before

QUE: VOTE_HOW_POST**ASK IF VOTED (VOTED=3):**

How did you vote in the election?

[PN: RANDOMIZE RESPONSE OPTIONS; INCLUDE RANDOMIZATION IN DATA FILE]

- 1 In person at a polling place
- 2 By absentee or mail-in ballot

QUE: VOTEINPWHEN**ASK IF VOTED IN PERSON (VOTE_HOW_POST=1):**

When did you vote?

- 1 Before Election Day
- 2 On Election Day

QUE: KNOWSITU**ASK ALL:**

What is one thing that you would like politicians in Washington, D.C. to know about your own situation when they are writing laws and setting policy? Please share as much detail as you can.

[PN: OPEN END TEXT BOX]

[PN: ALLOW RESPONDENT TO ENTER AS MUCH TEXT AS THEY WOULD LIKE WITH NO CHARACTER LIMIT]

BAT: THERMO**ASK ALL:**

We'd like to get your feelings toward a number of people on a "feeling thermometer." A rating of zero degrees means you feel as cold and negative as possible. A rating of 100 degrees means you feel as warm and positive as possible. You would rate the person at 50 degrees if you don't feel particularly positive or negative toward them.

[PN: RANDOMIZE ITEMS]

[Enter the number in the box between 0 and 100 that reflects your feelings]

BATTERY ITEMS:

THERMTRUMP How do you feel toward Donald Trump?

THERMBIDEN How do you feel toward Joe Biden?

THERMHARRIS How do you feel toward Kamala Harris?

QUE: VTCOUNT_OWN**ASK IF VOTED (VOTED=3):**

How confident are you that your vote was accurately counted?

- 1 Very confident
- 2 Somewhat confident
- 3 Not too confident
- 4 Not at all confident

[PN: RANDOMIZE ORDER OF VTCOUNT_POST_INP AND VTCOUNT_POST_ABS]**QUE: VTCOUNT_POST_INP****ASK ALL:**

How confident are you that votes cast in person at polling places across the United States were counted as voters intended in the elections this November?

- 1 Very confident
- 2 Somewhat confident
- 3 Not too confident
- 4 Not at all confident

QUE: VTCOUNT_POST_ABS**ASK ALL:**

How confident are you that votes cast by absentee or mail-in ballot across the United States were counted as voters intended in the elections this November?

- 1 Very confident
- 2 Somewhat confident
- 3 Not too confident
- 4 Not at all confident

QUE: VISITISS**ASK ALL:**

Have you ever visited the International Space Station?

- 1 Yes
- 2 No

[PN: RANDOMIZE ORDER OF OPENIDEN AND GAINS]

QUE: OPENIDEN**ASK ALL:**

Please choose the statement that comes closer to your own views – even if neither is exactly right.

[PN: RANDOMIZE STATEMENTS]

- 1 America's openness to people from all over the world is essential to who we are as a nation
- 2 If America is too open to people from all over the world, we risk losing our identity as a nation

QUE: GAINS**ASK ALL:**

Please choose the statement that comes closer to your own views – even if neither is exactly right.

[PN: RANDOMIZE STATEMENTS]

- 1 The gains women have made in society have come at the expense of men
- 2 The gains women have made in society have not come at the expense of men

Note: Response options for this question were not randomized due to a programming error.

QUE: DIVERSEPOP_MOD**ASK ALL:**

In general, do you think the fact that the U.S. population is made up of people of many different races, ethnicities and religions... **[PN: RANDOMIZE 1 AND 2, 3 ALWAYS LAST]**

- 1 Strengthens American society
- 2 Weakens American society
- 3 Doesn't make much difference **[anchor]**

QUE: TRANSGEND1**ASK ALL:**

Which statement comes closer to your views, even if neither is exactly right?

[PN: RANDOMIZE RESPONSE OPTIONS]

- 1 Whether someone is a man or a woman is determined by the sex they were assigned at birth
- 2 Someone can be a man or a woman even if that is different from the sex they were assigned at birth

QUE: RGHTCNTRL**ASK ALL:**

What do you think is more important?

[PN: RANDOMIZE]

- 1 To protect the right of Americans to own guns
- 2 To control gun ownership

BAT: SOCIETY**ASK ALL:**

Do you think each of the following is generally good or bad for our society?

[PN: RANDOMIZE ITEMS. RANDOMIZE RESPONSE OPTIONS 1-5 OR 5-1]

BATTERY ITEMS:

SOCIETYSSM	Same-sex marriages being legal in the U.S.
SOCIETYWHT	White people declining as a share of the U.S. population
SOCIETYEC	More people using electric vehicles instead of gasoline vehicles
SOCIETYBC	Birth control pills, condoms and other forms of contraception being widely available
SOCIETYTHER	More people openly discussing mental health and well-being
SOCIETYGUNS	An increase in the number of guns in the U.S.

RESPONSE OPTIONS:

- 1 Very good for society
- 2 Somewhat good for society
- 3 Neither good nor bad for society
- 4 Somewhat bad for society
- 5 Very bad for society

QUE: ABORTLGL**ASK ALL:**

Do you think abortion should be...

[PN: RANDOMIZE DISPLAY OF OPTIONS 1-4 AND 4-1]

- 1 Legal in all cases
- 2 Legal in most cases
- 3 Illegal in most cases
- 3 Illegal in all cases
- 4

QUE: RACESURV17**ASK ALL:**

How much, if at all, do you think the legacy of slavery affects the position of Black people in American society today?

- 1 A great deal
- 2 A fair amount
- 3 Not much
- 4 Not at all

QUE: CLMSRS1**ASK ALL:**

In your view, is global climate change a...

- 1 Extremely serious problem
- 2 Very serious problem
- 3 Somewhat serious problem
- 4 Not too serious problem
- 5 Not a problem

QUE: LGLSTATUS**ASK ALL:**

Which comes closer to your view about how to handle undocumented immigrants who are now living in the U.S.?

[PN: RANDOMIZE]

- 1 They should not be allowed to stay in the country legally
- 2 There should be a way for them to stay in the country legally, if certain requirements are met

QUE: WHADVANT**ASK ALL:**

In general, how much do White people benefit from advantages in society that Black people do not have?

- 1 A great deal
- 2 A fair amount
- 3 Not too much
- 4 Not at all

QUE: VOTED2020**ASK ALL:**

Did you happen to vote in the 2020 presidential election?

- 1 Yes
- 2 No

QUE: VOTEGEN2020**ASK IF VOTED (VOTED2020=1)**

In the 2020 presidential election, who did you vote for?

[PN: RANDOMLY DISPLAY RESPONSE OPTIONS 1-2 OR 2-1 FIRST]

- 1 Donald Trump, the Republican
- 2 Joe Biden, the Democrat
- 3 Jo Jorgensen, the Libertarian Party candidate
- 4 Howie Hawkins, the Green Party candidate
- 5 Another candidate

QUE: RELIG**ASK ALL:**

What is your present religion, if any?

- 1 Protestant (for example, Baptist, Methodist, non-denominational, Lutheran, Presbyterian, Pentecostal, Episcopalian, Church of Christ, Congregational/United Church of Christ, Holiness, Reformed, Church of God, etc.)
- 2 Roman Catholic
- 3 Mormon (Church of Jesus Christ of Latter-day Saints or LDS)
- 4 Orthodox (such as Greek, Russian, or some other Orthodox church)
- 5 Jewish
- 6 Muslim
- 7 Buddhist
- 8 Hindu
- 9 Atheist
- 10 Agnostic
- 11 Something else (please specify): [PN: INSERT SINGLE LINE TEXT BOX]
- 12 Nothing in particular

QUE: CHR**ASK IF SOMETHING ELSE, REFUSED/WEB BLANK TO RELIG (RELIG=11,99):**

Do you think of yourself as a Christian?

- 1 Yes
- 2 No

QUE: BORN**ASK IF CHRISTIAN (RELIG=1-4 OR CHR=1):**

Would you describe yourself as a born-again or evangelical Christian?

- 1 Yes, born-again or evangelical Christian
- 2 No, not born-again or evangelical Christian

QUE: PARTY**ASK ALL:**

In politics today, do you consider yourself a...

~~[PN: RANDOMIZE OPTIONS 1-2]~~

- 1 Republican
- 2 Democrat
- 3 Independent
- 4 Something else

Note: Response options for this question were not randomized due to a programming error.

QUE: PARTYLN**ASK IF INDEP, SOMETHING ELSE, OR REFUSED (PARTY=3, 4, 99):**

As of today, do you lean more to...

[PN: RANDOMIZE OPTIONS 1-2 IN SAME ORDER AS PARTY]

- 1 The Republican Party
- 2 The Democratic Party

QUE: IDEO**ASK ALL:**

In general, would you describe your political views as...

[PN: ROTATE OPTIONS 1-5/5-1]

- 1 Very conservative
- 2 Conservative
- 3 Moderate
- 4 Liberal
- 5 Very liberal

QUE: REG**ASK ALL:**

Which of these statements best describes you?

- 1 You are absolutely certain that you are registered to vote at your current address
- 2 You are probably registered, but there is a chance your registration has lapsed
- 3 You are not registered to vote at your current address

QUE: BIRTHPLACE**ASK ALL:**

Where were you born?

- 1 U.S. – 50 states, District of Columbia
- 2 U.S. – Puerto Rico
- 3 U.S. – other territory
- 4 Another country (please specify): **[PN: INSERT SINGLE LINE TEXT BOX]**

QUE: YEARSINUS**ASK IF BORN IN PUERTO RICO, OTHER US TERRITORY OR OUTSIDE OF U.S.****(BIRTHPLACE=2,3,4):**

How many years have you lived in the United States (excluding Puerto Rico or other U.S. territories)?

[Enter 0 for less than one year.]

— [PN: NUMERIC TEXT BOX]

QUE: ICEBREAKER**ASK ALL:**

Have you ever served on a Polar-class icebreaker ship?

- 1 Yes
- 2 No

QUE: PEAFEVER_CPS**ASK ALL:**

Did you ever serve on active duty in the U.S. Armed Forces?

- 1 Yes
- 2 No

QUE: VOL12_CPS**ASK ALL:**

In the past 12 months, did you spend any time volunteering for any organization or association? (This includes activities people may not think of, such as infrequent activities or for children's schools)

- 1 Yes
- 2 No

QUE: CITIZEN**ASK ALL:**

Are you a citizen of the United States?

- 1 Yes
- 2 No

QUE: PIICONSENT**ASK ALL:**

Thank you for answering our questions! This study has been sponsored by Pew Research Center. To improve the quality of our research, we would like to combine your answers to this survey with information from a national voting registry. To do this, we need your name and home address. Providing this information is strictly voluntary. This information will be used for research purposes only. It will not be made public, sold, or used to contact you, and will be deleted after the study is complete.

You can learn out more about the Pew Research Center [here](#). You can find out how we will protect your privacy [here](#).

Are you willing to share your name and home address?

- 1 Yes
- 2 No

QUE: PIICOLLECT**ASK IF YES IN PIICONSENT (PIICONSENT = 1)**

Please enter your name and home address below.

First name: _____

Last name: _____

Address 1: _____

Address 2: _____

City: _____

State: _____

Zip: _____