

Political Communication Report
Summer 2026 - Issue 33
“Political Communication Research under Data Constraints”

The Price of Convenience: Commercial Panels and Survey Data Quality in Political Communication¹

Pablo Jost, Johannes Gutenberg-Universität, Mainz

Survey-based research has long occupied a central position in political communication scholarship. The dominant paradigm of measuring attitudes, tracking media use, and testing the effects of political messages rests heavily on self-administered questionnaires, and the volume of studies relying on them has grown substantially over the past two decades (Rains et al., 2018). That growth and the spread of commercial online access panels are not unrelated. Offering rapid fielding, geographic reach, and comparatively low costs, these convenience panels removed what had previously been a meaningful barrier to large-scale data collection (Hays et al., 2015). The result was a genuine expansion of empirical capacity. The author of this essay is no exception. Commercial panels have featured in my own work, and that complicity warrants acknowledgment before the critique begins.

The methodological debate accompanying this expansion is not new. Concerns about representativeness, selection bias, and response validity have been documented since online panels first entered mainstream social science practice (Baker et al., 2010; Berry et al., 2022). What has received less systematic attention is a more fundamental reframing of what panel choice involves. Recruitment is a design choice that fundamentally shapes who responds, how they engage, and what the data are capable of telling us — and what not.

This matters for social science broadly, but carries particular weight in political communication research. Trust in political institutions, credibility judgments of news, exposure to political messages, and political knowledge are not arbitrary outcomes. They carry direct democratic stakes, and they are also variables on which survey respondents may differ systematically depending on how they were recruited. In the following, I develop this argument across the dimensions of survey quality, especially concerning commercial panels, before turning to what

¹ Copyright © 2026 Pablo Jost. Licensed under the Creative Commons Attribution Non-commercial No Derivatives (by-nc-nd). Available at <http://politicalcommunication.org>.

a more reflective recruitment practice might look like. The argument is not a new one, and that is partly the point.

The Quality Costs of Convenience

The quality costs of commercial panels — convenience samples by design — are not uniform. Different dimensions carry different risks, and which risks matter most depends on what a study is trying to establish.

Composition

The most fundamental concern is who ends up in a commercial panel in the first place. Nonprobability recruitment through self-selection, banner advertising, and referral networks produces samples that deviate from population benchmarks in ways that standard demographic weighting cannot fully correct (Baker et al., 2010; Callegaro et al., 2014). The critical point for political communication research is that these deviations are not incidental to the constructs we study. Researchers using commercial panels to estimate levels of institutional trust, political cynicism, or media credibility risk conflating recruitment artifacts with the phenomena themselves. People who join incentivized panels for small payments may differ systematically from the general population on precisely those variables. Panel samples have been shown to skew toward lower openness to experience and more conservative political preferences relative to face-to-face counterparts (Valentino et al., 2020), though the pattern is not consistent across all comparisons (Hohenberg et al., 2025). The evidence is mixed, but that is itself a reason for caution: if such differences exist and go unmeasured, the data do not describe citizens but a particular type of survey participant.

Cross-national comparisons can compound these problems further. Online population composition varies structurally across countries, and those differences translate into representativeness problems that standard controls cannot resolve. What appears as a national difference in panel data may reflect differential panel composition rather than genuine contextual variation (Maslovskaya & Lugtig, 2022).

Measurement quality

Beyond who responds, there is the question of how they respond. Respondents who participate routinely and for financial reward develop low-effort response strategies: endorsing midpoints, selecting the same option across items, or completing questionnaires at speeds inconsistent with genuine engagement (Krosnick, 1991). These behaviors are structurally encouraged by the incentive model. Panel members motivated primarily by financial reward invest less in response quality than those who participate out of genuine interest, a difference directly reflected in higher straightlining rates in convenience panels than in probability-based

counterparts (Cornesse & Blom, 2023). Both effects tend to increase with panel tenure as survey-taking becomes routine, producing measurement error.

A more recent and structurally related concern deserves brief mention. The financial incentive model that defines commercial panels creates conditions not only for low-effort responding but for outright substitution: AI agents can complete surveys at a fraction of the cost of human participation, and partial AI assistance on open-ended tasks is already widespread among crowdsourcing platform users (Veselovsky et al., 2023; S. Zhang et al., 2025), which current detection methods are effectively obsolete (Westwood, 2025).

Stimulus processing in effect research

This dimension is of particular interest for political communication research, where much of the empirical efforts rest on survey experiments: respondents are exposed to stimuli such as political messages or news articles and evaluate candidates or form credibility judgments afterwards. The theoretical logic of these designs assumes that respondents attend to the stimulus, process its content, and respond on that basis. When respondents rush through a questionnaire in a commercial panel environment, that logic is undermined. Westwood et al. (2022) have shown that inattentive responding contributed to inflated estimates of support for political violence in survey experiments, with meaningful consequences for how scholars, politicians, and journalists read public opinion on a politically sensitive topic.

Inferential challenges: two risks

The inferential consequences of low-quality panel data run in two directions. On one side, random measurement error introduced by careless responding attenuates observed associations toward zero, increasing the risk of false negatives: real relationships go undetected, or their magnitude is systematically underestimated (Callegaro et al., 2014). On the other side, the lower costs of commercial panel data have substantially increased the number of relationships researchers test within a single dataset and in general. When many tests are conducted and only significant results are reported, the false positive rate inflates, regardless of actual effect sizes (Franco et al., 2015). These two risks compound: a field that simultaneously underestimates some effects and overreports others based on low-quality data is not converging on accurate knowledge.

Towards More Reflective Recruitment Practices

The problems described above do not affect all research designs equally, and the appropriate response depends on what a study is trying to establish.

The first step should be to search existing infrastructures before planning new data collection. High-quality probability-based datasets such as the European Social Survey, the American National Election Study, or the German Longitudinal Election Study are underused in political communication research relative to their potential, and some — including the GLES — periodically issue calls for questions that allow researchers to add items to forthcoming waves. Experimental research is somewhat more forgiving of convenience samples, since randomization distributes compositional characteristics across conditions. Student samples, snowball designs, and academic opt-in panels such as the German SoSci Panel are defensible alternatives for this purpose, but only when the researcher has explicitly considered whether sample characteristics interact with the mechanism under study (Leiner, 2019). A study on how message complexity affects political information processing carries a different risk when conducted on students than a study testing the relative effect of source framing on attitude change: in the former, educational background and cognitive engagement are directly implicated in the theoretical claim; in the latter, they are not. That distinction matters and should be made explicit in the research design rather than resolved by convention. It is also worth noting that the assumption of commercial panels outperforming student samples on quality grounds does not hold empirically: direct comparisons show that student samples collected in controlled laboratory settings match or outperform professional panels on multiple quality indicators, with commercial panels consistently showing the lowest overall response quality (Kees et al., 2017).

Social media recruitment occupies a distinct niche. It is appropriate both for experimental designs and for research whose subject is platform-specific behavior, where a sample drawn from the platform's actual user population is analytically motivated rather than merely convenient. The compositional risks are real nonetheless: who sees a survey invitation is shaped by platform logic rather than sampling design, whether invitations reach respondents through organic algorithmic distribution or paid advertising targeting. Open-access links are a primary vector for bot infiltration and coordinated manipulation (Höhne et al., 2025; Westwood, 2025). Social media recruitment seems defensible when the target population is genuinely platform-specific, or, for experimental designs, when there is no strong reason to expect that the platform's user population differs from a broader sample on the attributes relevant to the research question. In both cases, distribution should be as controllable as the platform allows, with open-access public links avoided where possible, and traffic monitoring combined with quality checks beyond standard attention items should be applied.

Regardless of recruitment mode, transparent reporting of data quality indicators should be a non-negotiable baseline. The field's current practice of noting attention check pass rates and proceeding is insufficient in two respects. First, standard instructional manipulation checks and straightlining indicators are of limited value in commercial panel environments, where experienced panelists have encountered them repeatedly and can pass them without genuinely engaging (Anduiza & Galais, 2016; Schonlau & Toepoel, 2015). Second, quality checks should

be appropriate to the research design: for studies relying on stimulus exposure, time-on-page measures and recall or recognition questions administered after the stimulus are more informative than generic attention checks, since they directly test whether the central assumption of the design holds (C. Zhang & Conrad, 2014). Whatever checks are used, the decisions that follow from them should be transparent. Excluding flagged respondents is itself an analytical choice that can widen pre-existing sample biases rather than correct them, and should be reported and justified rather than silently applied. Journals are well positioned to enforce these standards by treating quality indicator reporting as a methodological requirement rather than an optional supplement.

For research designs that genuinely require probability-based sampling, the cost barrier is real: obtaining a probability sample of modest size is often not in the budget (Kees et al., 2017). One place individual researchers can act immediately are grant proposals. Committing to a recruitment standard that fits the research purpose, and justifying that choice explicitly, makes a different kind of claim on funding agencies than simply minimizing costs. Beyond individual efforts, pooling resources across research groups to fund shared data collection and building multi-project designs that allow direct comparison across recruitment modes are steps that remain underexplored. The field would also benefit from treating existing high-quality datasets as a shared resource. Many surveys fielded with probability-based methods are never fully analyzed, and making these data more accessible, particularly to early career researchers who cannot bear the costs of fielding their own studies, would partially address the resource asymmetry that pushes less-resourced scholars toward commercial panels by default. Beyond the matter of fairness, there is also a more self-interested argument for sharing. In fields where data sharing is established practice, well-documented datasets accumulate citations independently of whatever the original team publishes from them, which gives researchers a direct professional incentive to share rather than archive their data. Political communication has not yet built that incentive structure, but there is no reason it could not.

The Harder Question

The alternatives outlined in the preceding section are not new. Probability-based panels, harmonized surveys, student laboratory samples, and open calls for questions have all been available to political communication researchers for years, some for decades. The methodological critique of commercial panels is equally well established. If the problem were simply one of awareness, the field would have acted on it by now.

The more plausible explanation is structural. Awareness of the problem is not what the field lacks. Commercial panels persist because the incentive architecture of academic publishing makes them rational. A study fielded on a commercial panel can be collected, processed, and submitted within weeks. A collaborative design drawing on probability-based recruitment may take years to yield a publishable paper. In an environment that rewards output volume and

novelty of findings over methodological investment, the commercial panel is not a compromise so much as an adaptation. This is not comfortable to observe from a position of having made the same adaptation. Guilty as charged.

The “it depends” logic developed in the previous section is analytically valid. In practice, however, the sequence is often reversed. The commercial panel tends to be chosen first, on logistical or budgetary grounds, and the differentiation by purpose invoked afterward to justify a decision that was never actually made on those grounds. That reversal is the risk: the logic becomes a post hoc rationale rather than a genuine design principle. The distinction is worth preserving. It is a question that belongs at the design stage, not in the limitations section of a paper written after the data have already been collected. Occasionally, the honest answer at that stage is not to proceed.

What follows from this is that methodological improvement at the level of individual recruitment decisions is necessary but not sufficient. The persistence of default-mode usage of commercial panels reflects a collective action problem: the field as a whole would benefit from higher standards, but no individual researcher can unilaterally adopt them without incurring real competitive costs. In a field whose central constructs carry stakes that extend beyond academic output, that problem is worth naming directly.

This asks something of researchers, but it asks something of those who evaluate their work too. Reviewers of manuscripts and grant proposals are in a position to demand clearer justification of recruitment decisions and fuller reporting of quality indicators. They are also in a position to resist treating an unfamiliar sample design as a weakness when it is, in fact, a justified choice that fits the research purpose.

Much of what this essay has argued for already exists. Probability-based infrastructures, calls for questions, and unanalyzed datasets sitting in repositories are resources the field has built and continues to underuse. Treating data constraints as a problem to be solved entirely through new collection, rather than through better use of what is already collected, gets the diagnosis only half right. None of this resolves the deeper incentive problem this essay has tried to describe. It does mean, though, that some of what reflective recruitment requires is already within reach, waiting less for funding than for the habit of looking.

References

- Anduiza, E., & Galais, C. (2016). Answering Without Reading: IMCs and Strong Satisficing in Online Surveys. *International Journal of Public Opinion Research*, edw007. <https://doi.org/10.1093/ijpor/edw007>
- Baker, R., Blumberg, S. J., Brick, J. M., Couper, M. P., Courtright, M., Dennis, J. M., Dillman, D., Frankel, M. R., Garland, P., Groves, R. M., Kennedy, C., Krosnick, J., Lavrakas, P. J., Lee, S., Link, M., Piekarski, L., Rao, K., Thomas, R. K., & Zahs, D. (2010). Research Synthesis: AAPOR Report on Online Panels. *Public Opinion Quarterly*, 74(4), 711–781. <https://doi.org/10.1093/poq/nfq048>
- Berry, C., Kees, J., & Burton, S. (2022). Drivers of Data Quality in Advertising Research: Differences across MTurk and Professional Panel Samples. *Journal of Advertising*, 51(4), 515–529. <https://doi.org/10.1080/00913367.2022.2079026>
- Callegaro, M., Villar, A., Yeager, D., & Krosnick, J. A. (2014). A critical review of studies investigating the quality of data obtained with online panels based on probability and nonprobability samples¹. In M. Callegaro, R. Baker, J. Bethlehem, A. S. Göritz, J. A. Krosnick, & P. J. Lavrakas (Eds.), *Online Panel Research* (1st ed., pp. 23–53). Wiley. <https://doi.org/10.1002/9781118763520.ch2>
- Cornesse, C., & Blom, A. G. (2023). Response Quality in Nonprobability and Probability-based Online Panels. *Sociological Methods & Research*, 52(2), 879–908. <https://doi.org/10.1177/0049124120914940>
- Franco, A., Malhotra, N., & Simonovits, G. (2015). Underreporting in Political Science Survey Experiments: Comparing Questionnaires to Published Results. *Political Analysis*, 23(2), 306–312. <https://doi.org/10.1093/pan/mpv006>
- Hays, R. D., Liu, H., & Kapteyn, A. (2015). Use of Internet panels to conduct surveys. *Behavior Research Methods*, 47(3), 685–690. <https://doi.org/10.3758/s13428-015-0617-9>
- Hohenberg, B. C. von, Ventura, T., Nagler, J., Menchen-Trevino, E., & Wojcieszak, M. (2025). Survey Professionalism: New Evidence from Web Browsing Data. *Political Analysis*, 1–19. <https://doi.org/10.1017/pan.2025.10018>
- Höhne, J. K., Claassen, J., Shahania, S., & Broneske, D. (2025). Bots in web survey interviews: A showcase. *International Journal of Market Research*, 67(1), 3–12. <https://doi.org/10.1177/14707853241297009>
- Kees, J., Berry, C., Burton, S., & Sheehan, K. (2017). An Analysis of Data Quality: Professional Panels, Student Subject Pools, and Amazon’s Mechanical Turk. *Journal of Advertising*, 46(1), 141–155. <https://doi.org/10.1080/00913367.2016.1269304>
- Krosnick, J. A. (1991). Response strategies for coping with the cognitive demands of attitude measures in surveys. *Applied Cognitive Psychology*, 5(3), 213–236. <https://doi.org/10.1002/acp.2350050305>
- Leiner, D. J. (2019). Too Fast, too Straight, too Weird: Non-Reactive Indicators for Meaningless Data in Internet Surveys. *Survey Research Methods*, 13(3), 229–248. <https://doi.org/10.18148/srm/2018.v13i3.7403>
- Maslovskaya, O., & Lugtig, P. (2022). Representativeness in Six Waves of Cross-National Online Survey (CRONOS) Panel. *Journal of the Royal Statistical Society Series A: Statistics in Society*, 185(3), 851–871. <https://doi.org/10.1111/rssa.12801>
- Rains, S. A., Levine, T. R., & Weber, R. (2018). Sixty Years of Quantitative Communication Research Summarized: Lessons from 149 Meta-Analyses. *Annals of the International*

- Communication Association*, 42(2), 105–124.
<https://doi.org/10.1080/23808985.2018.1446350>
- Schonlau, M., & Toepoel, V. (2015). Straightlining in Web survey panels over time. *Survey Research Methods*, 9(2), 125–137. <https://doi.org/10.18148/srm/2015.v9i2.6128>
- Valentino, N. A., Zhirkov, K., Hillygus, D. S., & Guay, B. (2020). The Consequences of Personality Biases in Online Panels for Measuring Public Opinion. *Public Opinion Quarterly*, 84(2), 446–468. <https://doi.org/10.1093/poq/nfaa026>
- Veselovsky, V., Ribeiro, M. H., & West, R. (2023). *Artificial Artificial Artificial Intelligence: Crowd Workers Widely Use Large Language Models for Text Production Tasks* (arXiv:2306.07899). arXiv. <https://doi.org/10.48550/arXiv.2306.07899>
- Westwood, S. J. (2025). The potential existential threat of large language models to online survey research. *Proceedings of the National Academy of Sciences*, 122(47), e2518075122. <https://doi.org/10.1073/pnas.2518075122>
- Westwood, S. J., Grimmer, J., Tyler, M., & Nall, C. (2022). Current research overstates American support for political violence. *Proceedings of the National Academy of Sciences*, 119(12), e2116870119. <https://doi.org/10.1073/pnas.2116870119>
- Zhang, C., & Conrad, F. (2014). Speeding in Web Surveys: The tendency to answer very fast and its association with straightlining. *Survey Research Methods*, 8(2), 127–135. <https://doi.org/10.18148/srm/2014.v8i2.5453>
- Zhang, S., Xu, J., & Alvero, A. (2025). Generative AI Meets Open-Ended Survey Responses: Research Participant Use of AI and Homogenization. *Sociological Methods & Research*, 54(3), 1197–1242. <https://doi.org/10.1177/00491241251327130>



[Pablo Jost](#) is a postdoctoral researcher at the Department of Communication at Johannes Gutenberg-University Mainz. His research is on political communication in the context of digitalization. Recent projects are focused on how (non-institutional) political and societal actors use digital platforms and how they adapt to the changing communication environment.