

Department of Economics

Partisan Bias in Economic Forecasts: Experimental Evidence from the Presidential Elections in Argentina

Diego Marino Fages

Working Paper No. 6, 2026

Department Economics

Durham University Business School

Mill Hill Lane

Durham DH1 3LB, UK

Tel: +44 (0)191 3345200

<https://www.durham.ac.uk/business/about/departments/economics/>

© Durham University, 2026

Partisan Bias in Economic Forecasts: Experimental Evidence from the Presidential Elections in Argentina*

DIEGO MARINO FAGES[†]

Abstract

Rising political polarization may affect a key input for policy: household economic expectations. I study whether political preferences shape forecasts of inflation, unemployment, and exchange rates in an online experiment during Argentina’s 2023 presidential election (N=1,162). The design exogenously varies accuracy incentives and information about current indicators. Accuracy incentives reduce gaps in candidate-contingent forecasts, especially for the black market exchange rate and unemployment. Information provision reduces forecast dispersion but has weaker effects on partisan gaps. The results matter for survey design when eliciting household expectations.

JEL: C9, D84, D91, E71

Keywords: Partisan Bias, Forecasts, Prediction, Expectations, Survey Experiment

* The experiment was approved under DUBS-2023-09-20T11_53_32-zqlr86 by the Ethics Committee at Durham University and pre-registered as “Motivated Forecasts: Elections in Argentina 2023 (#147198)” in www.aspredicted.org.

[†] Durham University, Mill Hill Lane, MHL 367, Durham, DH1 3LB, UK. E-mail: diegomarinofages@gmail.com. Personal Website: <https://sites.google.com/view/diegomarinofages/home>. I acknowledge the funds provided by Durham University Business School and thank Moira Carrio, Patricio Gonzalez, and Hernan Romero for helping to recruit participants. I also thank Olivier Armantier, Abigail Barr, Aleksei Chernulich, Julio Elias, Cristina Griffa, Seung-Keun Martinez, Patrick Maus, Agustina Sampaolesi, Andis Sofianos, Chris Starmer, Cole Williams, Marcelo Woo, and the members of the EMBR Research Centre at Durham, seminar participants from CEDEX and Nottingham, Newcastle, UCEMA, University College Dublin, SEET Malaga, NICEP 2024 Conference and IMEBESS Riga, SEEDEC Bergen, ESA World Meeting Bogota, ESA Helsinki, the CREE Experimental Workshop (Lima), and the Nordic Conference in Behavioral and Experimental Economics (Copenhagen) for useful feedback. I also gratefully acknowledge the hospitality of JILAEE while working on this project. Ronny Condor provided outstanding research assistance. Refine.ink was used to check the paper for consistency and clarity.

I. Introduction

Policymakers routinely rely on survey measures of expectations to assess how households perceive inflation, unemployment, exchange rates, and other key economic outcomes. Yet these measures are only useful if respondents report beliefs that reflect their information rather than their political loyalties (Bullock and Lenz, 2019). This concern has become increasingly important in an era of rising political polarization (Mian et al., 2023; Canen et al., 2020) and widespread misinformation (Vosoughi et al., 2018; Allcott and Gentzkow, 2017), where citizens may interpret even basic economic facts through a partisan lens. If political preferences shape stated expectations, survey-based measures may capture not only beliefs about the economy, but also expressive or motivated responses.

This paper studies whether individuals’ forecasts of economic indicators are systematically biased toward their preferred political candidates. Understanding this is important for both measurement and policy (Coibion et al., 2018). If survey expectations reflect partisan bias, then standard elicitation instruments may misstate how people actually perceive the economy, with consequences for the design of monetary and fiscal policy and for how policymakers communicate with the public (Coibion et al., 2022; D’Acunto et al., 2020).

To study this, I conduct an online experiment with a sample pool of students and alumni from one of Argentina’s largest public universities. The country held presidential elections on 22 October 2023, with polls conducted in the lead-up indicating a highly uncertain outcome. The two dominant parties, located at the extremes of the political spectrum, espoused divergent views on how to balance inflation and unemployment, and generally projected competing visions for what the government’s role in the economy should be. In addition, in September of 2023, Argentina faced monthly inflation of around 12%. This context provides an ideal setting to study how supporters of each party form beliefs about the future of the economy and how these beliefs might be subject to partisan bias.

In the experiment, I ask participants to forecast inflation, unemployment, and ex-

change rates for the following year under each possible electoral outcome. I provide monetary incentives to half of the participants. The idea is that, if these differences partly reflect partisan bias rather than genuine beliefs about candidates' performance, real stakes for accuracy should reduce the gap between forecasts across scenarios. This is because the incentive structure shifts the weight of the incentives towards accuracy (Kunda, 1990). In an orthogonal treatment, I provide participants with correct information about the current levels of the variables of interest before eliciting their forecasts (Coibion and Gorodnichenko, 2026).

The main result is that accuracy incentives reduce partisan forecast gaps. Incentives reduce the maximum gap across candidate-contingent forecasts for all four indicators, with the strongest and most precise effects for the blue dollar and unemployment. The blue-dollar result is especially informative because respondents are highly knowledgeable about its current value, suggesting that the effect is not driven only by ignorance about current conditions. Information provision has weaker and less systematic effects on partisan gaps, although it reduces forecast dispersion in some cases. This pattern suggests that inaccurate or heterogeneous starting points matter for precision, but do not fully explain the politically aligned component of forecasts.

Several features of the design help interpret these effects. Because respondents forecast the same indicators under different candidate scenarios, the analysis compares forecasts within respondent and is less likely to reflect differences in ability, optimism, or general economic knowledge. Moreover, this allows to disentangle whether the reduction in the gap is driven by a positive or negative bias. In this regard, I find that respondents tend to forecast better outcomes under their preferred candidate and worse outcomes under candidates they do not support. The bias decomposition shows that accuracy incentives and the information treatments mainly reduce pessimism about non-preferred candidates.

The heterogeneity analysis shows that the main findings are not driven by a single subgroup: accuracy incentives reduce forecast gaps across respondents with different levels of literacy, knowledge about the current values of the indicators, response time, and political leaning.

The paper’s main contribution is to provide causal evidence of the effect of using monetary incentives on reducing partisan bias on economic forecasts. To the best of my knowledge this is the first experiment incentivizing forecasts of real-world macroeconomic indicators.¹

The paper also contributes by eliciting and using conditional forecasts in the analysis.² Previous literature typically defines partisan gaps as differences between the forecasts of, for example, Democrats and Republicans (e.g., [Bachmann et al., 2021](#)). Even after controlling for many observable characteristics, these gaps remain confounded by unobserved differences across respondents. The methodological contribution of this paper is to compute forecast gaps as differences between the forecasts that the same individual makes under different candidate scenarios, holding fixed individual characteristics. These within-respondent gaps may partly reflect genuine beliefs about candidates’ economic consequences, but the experimental variation in accuracy incentives helps me to identify the politically motivated component of the gap.

Another strand of the literature focuses on eliciting expectations using surveys ([Dietrich et al., 2023](#); [Weber et al., 2022](#); [Armantier et al., 2013](#)). The current study suggests that the political orientation of the respondents needs to be taken into account to interpret their expectations, as individuals may not behave as they report in unincentivized surveys. On this strand, the study relates to a literature trying to understand why professional forecasters may not reveal their true beliefs ([Gemmi and Valchev, 2025](#); [Ottaviani and Sørensen, 2006a,b](#); [Laster et al., 1999](#); [Ehrbeck and Waldmann, 1996](#)). Although household incentives are less affected by profits than those of professional forecasters, I find that they are strongly affected by political motivations. These results suggest the need for a careful aggregation of household expectations, which are normally estimated as the mean or median ([Coibion et al., 2018](#)). This is especially important in countries like the US, where there is strong spatial segregation in terms of political leaning. This paper not only provides causal evidence of partisan but also

¹For examples using incentives for surveys regarding current economic conditions, see [Graham \(2025\)](#), [Bullock et al. \(2015\)](#), [Prior et al. \(2015\)](#).

²A small number of papers have already used conditional forecasts; however, none of them simultaneously introduces experimental variation either through monetary incentives or by providing current information about the indicators. For examples, see [Bauer et al. \(2023\)](#) and [Coibion et al. \(2020\)](#).

an effective (and inexpensive) method to significantly reduce them in surveys eliciting individuals' forecasts.

Partisan bias in economic expectations also has implications for survey design and policy communication. Elicitation instruments may need to account for political preferences when outcomes are politically salient, especially if unincentivized responses partly reflect expressive reporting. Similarly, if households process economic information through a political lens, policymakers may need to account for this in the content of their messages, the medium and messenger (Coibion et al., 2022; Mokhtarzadeh and Petersen, 2021; D'Acunto et al., 2020).

While most studies of partisan bias focus on beliefs about current or past events, I study forecasts of future, unrealized outcomes. The paper therefore relates to the literature on expectation formation and departures from full-information rational expectations (Coibion et al., 2018), including work on inflation experiences (Malmendier and Nagel, 2016), information provision (Cavallo et al., 2017; Armantier et al., 2016, 2013), identity (Bauer et al., 2023; Donkor et al., 2023), and demographic factors (see D'Acunto and Weber, 2024, for a survey). I contribute by studying political preferences as a driver of macroeconomic expectations.

Regarding the themes studied in the literature on partisan bias, many studies address general questions related to topics that are contentious (and emotional) for the respective political parties. It is possible that some of these questions are not relevant to the particular subjects, or at least not to the same extent. Hence, not everyone is equally informed about the topics (e.g., Obama's religion or the number of people in a photograph). It has been argued that subjects may employ heuristics that favor their political party when they are unsure about the answer (Bullock and Lenz, 2019). There may be greater reliance on these heuristics when the discussion concerns irrelevant topics (Bullock and Lenz, 2019). Furthermore, supporters of different parties might differ in the extent of their interest in certain questions. By contrast, in this study, I look at the economic indicators considered most important in the current Argentinean context, which are the same across candidates and are arguably high-stakes topics for everyone.

The study also contributes to the literature (mostly US based) by studying a sample

from a developing country with very unstable macroeconomic conditions (Argentina has one of the highest inflation rates in the world). In this study, forecast gaps respond modestly when respondents are provided with accurate information regarding the current values but the effect varies across indicators suggesting that even in these conditions some indicators are more salient than others. This is in line with the hypothesis that in more unstable economies, people are more attentive to macroeconomic variables (Weber et al., 2025; Cavallo et al., 2017).

The paper proceeds as follows. Section I describes the electoral and macroeconomic context. Section II presents the experimental design. Section III describes the sample and provides descriptive evidence. Section IV reports the main results, heterogeneity analyses, and robustness checks. Section V discusses the interpretation of the findings and concludes.

II. Context

In 2023, the primary elections produced three main candidates: Patricia Bullrich (Juntos por el Cambio), Sergio Massa (Union por la Patria), and Javier Milei (La libertad Avanza).³ Sergio Massa was the head of the Ministry of Economy and responsible for implementing strong populist measures leading to high inflation. Patricia Bullrich was expected to be the strongest opposition candidate, representing the party that won the elections in 2015. Javier Milei was the founder of a new party representing the far right, with very strong views about the appropriate role for the state in the economy (e.g., proposing to ‘blow up’ the Central Bank). The PASO elections (simultaneous and mandatory open primaries) were very close: La Libertad Avanza (30.04%), Union por la Patria (27.3%), and Juntos por el Cambio (28.3%).

The first round of the elections took place on 22 October 2023. The leading candidate in this round was Massa with 36.68%, followed by Milei with 29.98%, Bullrich with 23.82%, Schiaretti with 6.78%, and Bregman with 2.7%. The second round, held on 19 November, was won by Milei with 55.69% of the vote.

³The other two candidates were Juan Schiaretti (Hacemos por nuestro país) and Myriam Bregman (Frente de Izquierda).

III. Experimental Design

The experiment was conducted online with students and alumni from one of Argentina’s largest public universities (Universidad Nacional del Nordeste)(N=1162).⁴ The experiment was strategically conducted the week prior to the presidential elections (the survey was open from 16–21 October 2023).^{5 6} This is instrumental for the experimental design because I ask subjects to forecast economic variables for the first few months of the new presidency. I elicit these forecasts using the strategy method such that the forecasts are made conditional on each possible candidate winning the election.

All subjects in the experiment are asked the same questions, that is, regarding the current levels of inflation, unemployment, and exchange rates for the US dollar in the official and black (called blue in Argentina) market;⁷ and to make forecasts about the future levels of these variables for each possible electoral outcome. I also ask participants to forecast unrelated events, such as football outcomes in the national league, temperatures, and the change in the price of bitcoin and the Merval index on a specific day. The online Appendix contains the main screens, with the English translations (see the replication package for the full survey in Spanish and English).

I introduce two treatments in a factorial 2x2 design. For participation, all subjects are rewarded with a chance to win one of the thirty 100USD cash prizes. The decision to pay in US dollars was intended to prevent inflation from devaluing the prize and to make it more salient (people did not know how many participants were taking part in the survey). Half of the sample could obtain 10 extra lottery tickets for each accurate economic forecast (i.e., if their prediction fell within 5% below or above the actual value), and the other half could obtain 10 extra lottery tickets for each correct forecast

⁴The experiment was programmed in Qualtrics. It was approved under DUBS-2023-09-20T11_53_32-zqlr86 by the Ethics Committee at Durham University and pre-registered in AsPredicted (#147198). For convenience, it can also be found in the Online Appendix.

⁵Figure OA.3 in the Online Appendix shows the distribution of responses over time.

⁶Argentina hosts two presidential debates before the elections. In the first debate, on October 1st, the candidates discussed economics, education, human rights, and democracy. The second debate took place on October 8th, and security, unemployment, production, human development, housing, and environmental protection were the topics discussed.

⁷Argentina has a multiple exchange rate system. See <https://www.bloomberglinea.com/english/a-beginners-guide-to-understanding-argentinas-multiple-exchange-rates/>.

in other unrelated topics (i.e., football, weather and stock markets). I will refer to this as the accuracy treatment.⁸ The idea is that the treatment increases the cost of being wrong; hence, subjects' forecasts are less likely to be distorted by emotions, desires, and goals.⁹ ¹⁰ Furthermore, the direction of the correction is endogenous to the subjects' political identity (Bénabou and Tirole, 2016). The second, orthogonal treatment, is to inform participants about the current levels of the indicators after they have answered the questions regarding current levels but before they provide their forecasts (i.e., information treatment). I also do this for half of the sample. Figure 1 depicts the treatment arms with the four conditions.

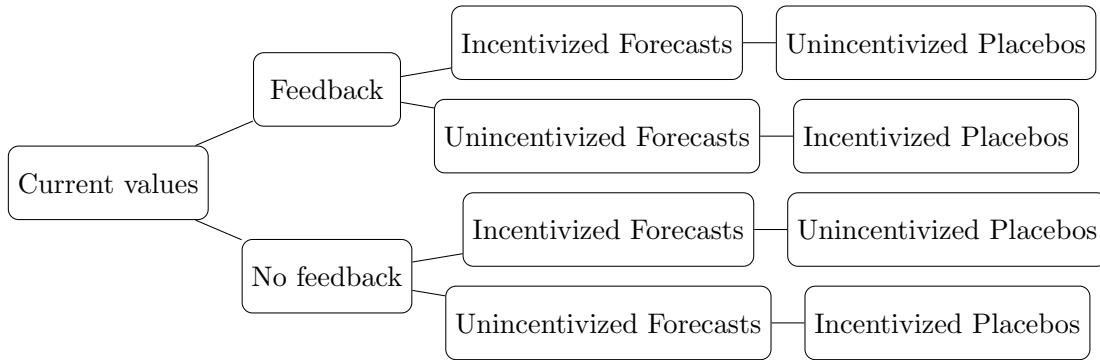


Figure 1 : Treatment arms

This design presents two main desirable features: First, since subjects need to forecast the future, there is no need to worry about participants knowing the correct answer, and they cannot cheat by Googling the answer (especially because of the highly uncertain economic and political context). Second, by asking about the future, we can ask questions conditional on specific outcomes, which allows us to observe differences in counterfactual scenarios. Although the effective stakes attached to each conditional forecast depend on the respondent's perceived probability that the corresponding can-

⁸Before the extra chances earned by correct responses, the expected value of the prize was $3000/1162 = 2.58\text{USD}$.

⁹Note that, for all participants—regardless of their preferred candidate or the election outcome—lower values across all indicators are more economically desirable.

¹⁰This approach is related to recent pay-for-correct designs that elicit beliefs about economic statistics before their official release (Graham, 2025).

didate will win, the incentives remain aligned with accuracy: for any candidate assigned positive probability, a more accurate conditional forecast weakly increases the expected chance of earning the bonus.

I then asked the participants about their preferred candidates, who they thought would win, who they would vote for in the first round, and some literacy questions.

IV. Sample and descriptive statistics

The subject pool consists of 1,162 students and alumni from the School of Economic Sciences at Universidad Nacional del Nordeste.¹¹ The university draws many students from Corrientes and Chaco, two provinces with different political traditions (Chaco being traditionally Peronist and Corrientes being traditionally Radical), and the region is one of the least developed in Argentina. The survey was fielded between October 16 and October 21, 2023, immediately before the first round of the presidential election.

Public universities in Argentina are free, and admission is not binding in practice. This generates high enrollment and high dropout rates in the early stages of study. The sample includes both students and alumni, whether or not they completed their degrees. Although the sample is not nationally representative, it is highly heterogeneous, especially politically. Among respondents who stated a favorite candidate, 55% preferred Milei; in the full sample, including nonresponses, this share is 49%, close to Milei’s vote share in the second round of the election.

Table OA.1 shows that the favorite-candidate measure is closely related to respondents’ intended vote. Among respondents who state a favorite candidate, 87.5% of Massa supporters, 89.2% of Bullrich supporters, and 92.2% of Milei supporters report intending to vote for that same candidate. The table also shows that favorite candidate is distinct from expectations about the winner: most respondents expected Milei to win, including 58.9% of Massa supporters and 69.2% of Bullrich supporters and 93.7% of Milei supporters. Preferences regarding the second best option also reveal the main political alignment in the sample: Bullrich and Milei supporters overwhelmingly list

¹¹This school includes the departments of accountancy, economics, and management. The metropolitan area where the university is located has approximately 727,000 inhabitants.

each other as their second preference, while Massa supporters mainly support Bullrich as second.

Table OA.2 reports balance across accuracy-treatment arms. Participants are, on average, 26 years old, 57% are women, and 58% are employed. On a scale from 0 to 10, the average agreement with right-wing policies is 5.7. The only statistically significant imbalance is a small difference in the share of women across treatment arms. The treatment and control groups are balanced in economic and financial literacy (Table OA.3) and in their reported current values of the economic indicators (Table OA.4). Figures OA.1, OA.2, and OA.3 show the distributions of survey completion dates, literacy scores, and political leaning.

Figure OA.4 plots reported current values by literacy score, separately by accuracy-treatment arm, and compares them with the benchmark values. Respondents are remarkably accurate about the blue dollar, the indicator that is most salient in everyday economic discussions and is reported daily in newspapers. For the other indicators, more literate respondents report values closer to the benchmarks, especially for the official dollar, inflation, and unemployment. The figure also highlights the large baseline misperception about unemployment: even among high-literacy respondents, reported unemployment remains far above the actual value.

I next examine whether reported current values vary with political preferences. Because Massa was the ruling coalition's candidate and Minister of Economy, partisan bias would imply that Massa supporters report better current economic conditions than opposition supporters. Table OA.5 provides evidence consistent with this pattern. Massa supporters report lower values for all four indicators.¹² The differences are statistically significant for unemployment and marginally significant for inflation. In levels, Massa supporters report an unemployment rate of 31%, compared with about 35% among Bullrich and Milei supporters (Table OA.6). They also report lower inflation than the other groups. Consistent with this pattern, respondents are very unlikely to identify their preferred candidate as the most corrupt candidate and in almost all

¹²Figure OA.5 compares the absolute percent deviation of reported current values from the benchmark values by candidate support.

cases expect the best outcomes for their own candidate (Table OA.7, Panel B). Thus, the descriptive evidence is suggestive of partisan bias, but it is not conclusive. Massa supporters report somewhat better current economic conditions under the incumbent government, but these differences could also reflect broader differences in information or attention across political groups that are in power versus the opposition. The next section therefore turns to candidate-contingent forecasts, where the same respondent reports expected outcomes under different possible electoral winners.

V. Results

A. Accuracy treatment

The main hypothesis is that accuracy incentives reduce the gap between respondents' candidate-contingent economic forecasts. To study this, Table 1 presents the results of regressing the gaps on the accuracy treatment while restricting the sample to respondents who did not receive the information treatment economic conditions. The first row reports the treatment effect on the maximum difference across the three candidate scenarios. The coefficients are negative for all four outcomes. The effect is particularly large for the blue dollar, where the maximum gap falls by 206 pesos, or 24% of the control-group mean, and for unemployment, where the gap falls by 4.9 percentage points, also about 24% of the control-group mean. The estimates for the official dollar and inflation are also negative, but less precisely estimated.

The remaining rows present pairwise comparisons between candidates. For example, the "Massa and Milei only" row is restricted to Massa and Milei supporters, and the gap is defined as $Forecast_{Milei} - Forecast_{Massa}$ for Massa supporters and as $Forecast_{Massa} - Forecast_{Milei}$ for Milei supporters. These comparisons show the same general pattern: most coefficients are negative, indicating that accuracy incentives reduce the extent to which respondents report worse forecasts under candidates they do not support. The last row reports the gap between the forecasts associated with the least and most preferred candidates.¹³ This estimate is also negative for all four

¹³This gap was not preregistered. The maximum-gap measure was preregistered as the main outcome because

Table 1—: Accuracy Incentives and Economic-Indicator Forecast Gaps

	(1) Official dollar	(2) Dollar blue	(3) Inflation	(4) Unemployment
Maximum difference	-26.12 (58.76)	-206.30** (82.47)	-18.23 (12.50)	-4.94*** (1.59)
Avg. Control	553.22 (42.88)	871.09 (69.83)	81.08 (9.79)	20.33 (1.25)
Observations	530	529	474	530
Massa and Milei only	-92.79 (84.51)	-131.13 (116.28)	-5.24 (16.70)	-4.53** (2.26)
Avg. Control	211.96 (58.29)	579.96 (98.32)	62.52 (12.87)	14.89 (1.81)
Observations	383	380	341	383
Massa and Bullrich only	-30.19 (61.61)	-119.10 (72.35)	-9.78* (5.17)	-2.21 (1.98)
Avg. Control	149.93 (52.03)	360.11 (48.92)	25.32 (3.96)	9.80 (1.64)
Observations	225	225	204	225
Milei and Bullrich only	-167.88*** (59.82)	-111.72 (83.40)	-15.29 (10.73)	-4.14*** (1.58)
Avg. Control	159.83 (45.22)	328.63 (73.17)	35.46 (9.95)	8.98 (1.27)
Observations	426	424	381	426
Least - most favourite	-91.58 (65.41)	-123.86 (86.63)	-7.56 (12.31)	-4.82*** (1.77)
Avg. Control	205.50 (46.13)	545.10 (71.08)	53.83 (9.22)	14.63 (1.42)
Observations	517	515	463	517

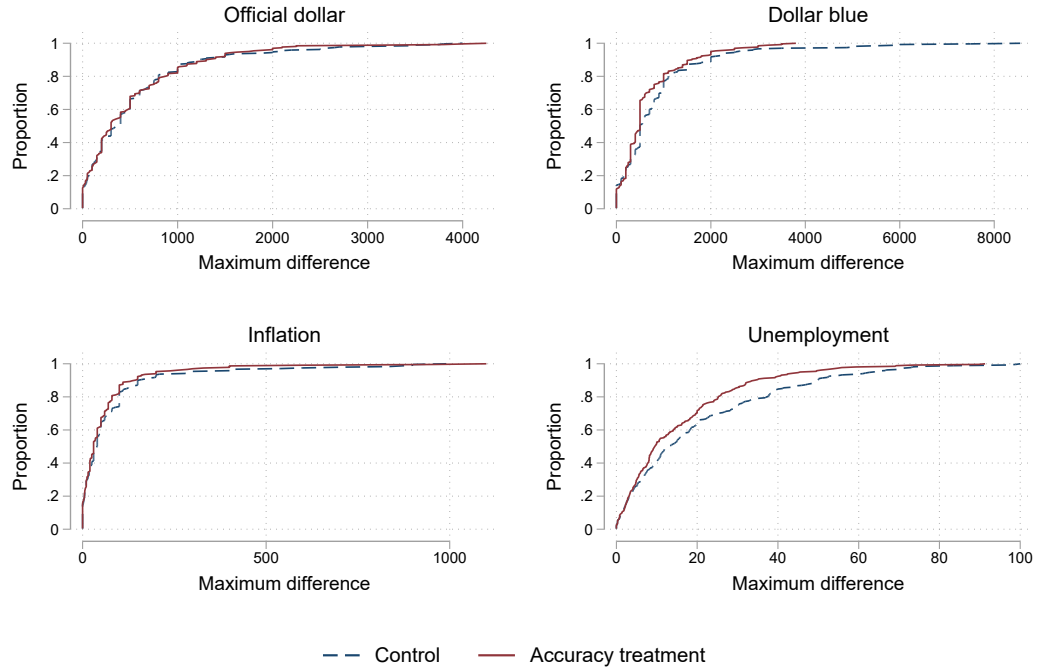
Notes: The table reports coefficients from regressions of forecast-gap measures on the accuracy-treatment indicator, restricted to respondents who did not receive the information treatment. The main coefficients are in the first row, where the outcome is the maximum difference across the three candidate-contingent forecasts for each indicator. The remaining rows report pairwise candidate gaps and the least-minus-most preferred candidate gap. Robust standard errors are in parentheses. The control-group mean is reported for each outcome. *, **, and *** indicate significance at the 10%, 5%, and 1% levels.

outcomes, and is statistically significant for unemployment. My preferred specification remains the maximum-gap measure, because it was preregistered and does not require respondents' candidate preferences to align mechanically with their expected economic performance.

it captures the overall dispersion in a respondent's candidate-contingent forecasts without imposing that the respondent's favorite and least-favorite candidates must define the most economically different scenarios. This is useful in a setting with three candidates and heterogeneous beliefs about candidate competence.

It is interesting to note that the strongest effects come from the two indicators for which respondents are the most and the least knowledgeable (i.e., people are constantly reminded about the blue dollar in their daily lives and very little informed about the unemployment rate).

Figure 2 : Distribution of Maximum Forecast Gaps by Accuracy-Incentive Status



Notes: The figure plots empirical cumulative distribution functions of the maximum forecast gap for each economic indicator, separately by accuracy-treatment status. The sample is restricted to respondents who did not receive the information treatment. The maximum forecast gap is defined as the difference between the highest and lowest candidate-contingent forecast across the three candidates. Corresponding Kolmogorov–Smirnov tests are reported in Table OA.9.

In order to see the full distribution of the gaps, Figure 2 presents the CDFs of the maximum forecast gaps by treatment status. Table OA.9 reports two-sample Kolmogorov–Smirnov tests for the maximum forecast gaps. In the comparison between accuracy incentives and the control group among respondents without information, the distribution of blue-dollar gaps differs significantly across treatment arms ($D = 0.152$, $p = 0.005$), as does the distribution of unemployment gaps ($D = 0.121$, $p = 0.041$). The corre-

sponding differences for the official dollar and inflation are not statistically significant. The same table shows that the information treatment, even without accuracy incentives, shifts the distribution of unemployment gaps, and that the combined treatment differs from the pure control group mainly for unemployment. Thus, the distributional evidence reinforces the main result for unemployment and shows an additional shift in blue-dollar gaps under accuracy incentives.

Table 2—: Treatment Intensity and Maximum Forecast Gaps

Panel A: Treatment intensity				
	(1) Official dollar	(2) Dollar blue	(3) Inflation	(4) Unemployment
Treatment intensity	-0.0008 (0.0037)	-0.0114*** (0.0037)	-0.0082*** (0.0030)	-0.0087*** (0.0028)
Avg. Treated	48.85 (1.82)	41.96 (1.68)	45.08 (2.11)	55.00 (1.81)
Observations	511	511	458	511
Panel B: Placebo intensity				
[1em] Placebo intensity	0.0086** (0.0036)	0.0073** (0.0034)	0.0058 (0.0039)	0.0023 (0.0031)
Avg. Treated	48.82 (1.82)	41.93 (1.68)	45.08 (2.11)	55.06 (1.81)
Observations	511	511	458	511

Notes: The table reports coefficients from regressions of maximum forecast gaps, expressed as a percentage of the midpoint between the highest and lowest candidate-contingent forecasts, on the treatment-intensity measures. Treatment intensity is proxied by the respondent's blue-dollar forecast under the candidate they expected to win; higher values imply a higher expected local-currency value of the dollar-denominated prize. The placebo intensity measure is the mean of the blue-dollar forecasts under the candidates the respondents did not expect to win. Robust standard errors are in parentheses. The samples are restricted to respondents who received accuracy incentives. All regressions control for the information treatment. *, **, and *** indicate significance at the 10%, 5%, and 1% levels.

Because the prize was denominated in US dollars, respondents may have perceived different local-currency stakes depending on their exchange-rate expectations. I therefore explore whether, among respondents who received accuracy incentives, forecast gaps are smaller when the dollar-denominated prize is perceived to be more valuable in pesos.¹⁴ Table 2 presents the regression results. The outcomes are maximum gaps

¹⁴This exploratory analysis was suggested during a seminar presentation ex post and was not pre-registered. Since the perceived-stakes proxy is based on respondents' own forecasts, the estimates should not be interpreted as causal effects of incentive size.

expressed as a percentage of the midpoint between the highest and lowest candidate-contingent forecasts.¹⁵ Treatment intensity is proxied by the respondent’s blue-dollar forecast under the candidate they expected to win. The estimated coefficients are negative for all outcomes and statistically significant for the blue dollar, inflation, and unemployment. A 100-peso increase in the perceived-stakes proxy, about 0.13 standard deviations, is associated with reductions of 1.14, 0.82, and 0.87 percentage points in the blue-dollar, inflation, and unemployment gaps.¹⁶ By contrast, the placebo-intensity measure, based on the candidates respondents did not expect to win, does not show the same pattern. This evidence is therefore suggestive, but consistent with accuracy incentives being more effective when they are perceived as more valuable.

Finally, Table OA.10 decomposes the reduction in partisan forecast gaps into changes in forecasts under the respondent’s favorite candidate and changes in forecasts under other candidates. All outcomes are standardized using the no-accuracy, no-information control group, and higher values correspond to worse economic outcomes. The accuracy treatment has little effect on forecasts under the respondent’s favorite candidate. Instead, it lowers forecasts under other candidates by 0.16 standard deviations and reduces the other-minus-favorite gap by 0.12 standard deviations, about 35% of the control-group gap. In other words, the accuracy treatment reduces partisan forecast gaps primarily because respondents become less pessimistic about the economic scenarios associated with candidates they do not support, not because they become more pessimistic about their own candidate.¹⁷ This is consistent with Rathje et al. (2023). Appendix Table OA.11 provides a complementary analysis by regressing the forecasts on the treatments with respondent and candidate fixed effects. The results show a large penalty for non-preferred candidates even after controlling for respondent and candidate fixed effects, confirming that the gaps are politically directional. The treatment

¹⁵This scaling is especially important for the blue-dollar outcome because the treatment-intensity proxy is itself based on blue-dollar forecasts. Without scaling, the analysis could mechanically relate higher perceived prize values to larger peso-denominated forecast gaps.

¹⁶This corresponds to a 10,000-peso increase in the perceived value of a USD 100 prize if won. A one standard deviation increase in the proxy is about 744 pesos.

¹⁷This directional pattern also helps mitigate the concern that respondents simply make safer or more moderate forecasts when monetary incentives are present: a generic hedging response would not necessarily predict that the adjustment is concentrated in forecasts for non-preferred candidates.

interactions are negative in the pooled specification, consistent with the decomposition results, although the indicator-specific estimates are less precise.

These findings qualify the literature on motivated reasoning and selective exposure. Prior work shows that voters are more likely to retain congenial information and avoid information that challenges their political views (Jerit and Barabas, 2012; Bauer et al., 2023; Robertson et al., 2023), and that belief revision is often concentrated among individuals with weaker priors (Coibion et al., 2020). I cannot observe whether participants search for information during the experiment, nor which sources they consult. However, the results show that monetary incentives move forecasts in a less partisan direction. If respondents search for information, they appear to use it in a way that tempers their prior political beliefs; if they do not search, the incentives still induce them to report forecasts that are less favorable to their preferred political narrative.

B. Forecast accuracy and beliefs about the direction of the error

Because Milei won the election, the Milei-contingent forecasts can be compared directly with realized outcomes. Panel A of Table OA.12 reports treatment effects on the absolute distance between respondents' Milei-contingent forecasts and the realized future value. There is little evidence that either treatment systematically improved the accuracy of forecasts for the official exchange rate, the blue dollar, or inflation. The clearest effect appears for unemployment: both the accuracy treatment and the information treatment reduce the distance between respondents' forecasts and the realized value, with the effect especially large and precisely estimated for the information treatment.

Panel B provides a more exploratory ex-post comparison using the average of the two non-Milei contingent forecasts, excluding the realized winner's scenario, to see if there was a general tendency to the realized values. Both treatments reduce these forecast distances. The effects are strongest and most precise for the official dollar. Both treatments also significantly reduce the distance for unemployment, and the information treatment additionally reduces the distance for inflation. By contrast, the blue-dollar effects are negative but imprecisely estimated.

Panels C and D examine whether the treatments moved forecasts toward the current values shown in the information treatment rather than toward future realized values. This is useful because improvements in apparent accuracy could partly reflect anchoring on contemporaneous values (either provided by the information treatment or search encouraged by the monetary incentives), or respondents independently searching for these values, rather than better beliefs about future outcomes. The evidence is consistent with some movement toward current values, especially for inflation and unemployment. However, the results do not indicate a mechanical anchoring effect across all outcomes: for the exchange-rate variables, the coefficients are generally imprecise.

Inspired by [Thaler \(2024\)](#), I also examine respondents' beliefs about the likely direction of their own forecast errors.¹⁸ This question was asked only for the blue-dollar forecasts. Table [OA.13](#) compares, within supporter groups, whether respondents are more likely to expect having overestimated under their preferred candidate than under other candidates. The pattern differs sharply by candidate support. In the control group, Milei supporters are 46 percentage points more likely to expect overestimation under Milei, meaning that if their preferred-candidate forecast is wrong, they expect the blue dollar to be even lower and therefore even more favorable to their preferred candidate. Bullrich supporters display the opposite pattern, while Massa supporters show little difference. These differences persist across treatment arms, although the magnitudes vary.

Table [OA.14](#) shows that respondents did not accurately anticipate the direction of their forecast errors. Across treatment arms, only 34–39% of respondents correctly predicted whether their forecast would be too high or too low, significantly below 50% in the pooled sample. Both treatments bring the share closer to 50% by 3–5 percentage points in the pooled sample.

Table [OA.15](#) provides a second diagnostic of beliefs about the direction of forecast errors. The table restricts attention to the respondent's preferred-candidate blue-dollar

¹⁸In Thaler's design, respondents evaluate the credibility of sources that tell them that the median of their belief distribution is too high or too low; because the signal is centered on the respondent's own median, a Bayesian respondent should not infer source quality from the direction of the message. Since I am not eliciting the median, the results should be interpreted with care.

forecast and asks whether the expected error points toward the most different of the other two candidate-contingent forecasts. The pattern again differs sharply by candidate support. In the control group, Massa and Bullrich supporters are above the 50% benchmark, with shares of 74% and 84%, respectively. Milei supporters are below the benchmark, at 41%. This heterogeneity persists across treatment arms: Massa and Bullrich supporters generally expect errors to move toward the most different forecast, while Milei supporters generally expect errors to move away from it. Thus, the table does not show a uniform awareness of partisan bias; rather, beliefs about likely errors differ strongly by candidate support.

C. Information treatment

An orthogonal information treatment provided respondents with the correct current values of the indicators after they reported their beliefs about current economic conditions and before they submitted their candidate-contingent forecasts. This treatment was intended to align respondents' starting points with accurate information. Unlike the accuracy treatment, however, the information treatment did not create monetary incentives to report less politically motivated forecasts. Thus, the hypothesis is that information alone should have weaker effects on partisan forecast gaps, but may reduce dispersion by giving respondents a common factual anchor.

Table 3 tests whether the information treatment reduces the gap in forecasts by restricting the sample to respondents who did not receive accuracy incentives. The effects on maximum forecast gaps mostly go in the same direction as those of the accuracy treatment, although they are smaller and less precisely estimated. The maximum gap falls for the blue dollar, inflation, and unemployment, while the coefficient for the official dollar is positive and statistically insignificant. The clearest effect is for unemployment: the maximum gap falls by 3.7 percentage points, or about 18% of the control-group mean. The estimated reductions for the blue dollar and inflation are also negative, but not statistically significant.

The pairwise candidate comparisons show a similar pattern. Most coefficients are negative, suggesting that information provision tends to reduce partisan forecast gaps,

Table 3—: Information Treatment and Economic-Indicator Forecast Gaps

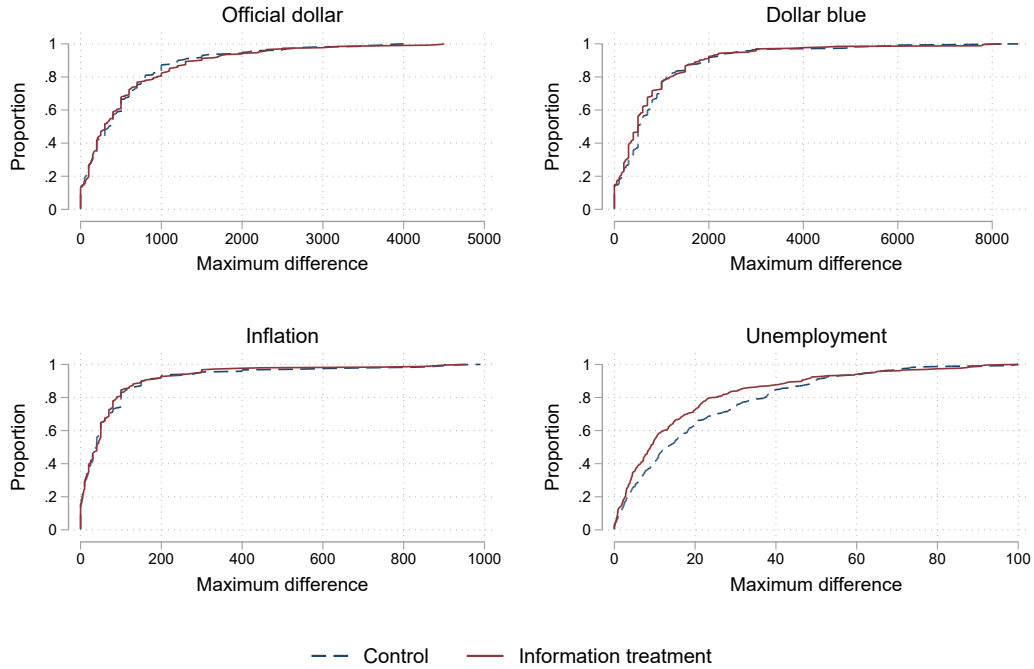
	(1) Official dollar	(2) Dollar blue	(3) Inflation	(4) Unemployment
Maximum difference	23.49 (63.72)	-52.83 (101.04)	-6.62 (13.11)	-3.74** (1.78)
Avg. Control	553.22 (42.88)	871.09 (69.83)	81.08 (9.79)	20.33 (1.25)
Observations	531	531	485	531
Massa and Milei only	-25.20 (88.50)	-52.99 (143.72)	2.14 (17.80)	-5.25** (2.60)
Avg. Control	211.96 (58.29)	579.96 (98.32)	62.52 (12.87)	14.89 (1.81)
Observations	361	361	326	361
Massa and Bullrich only	-40.24 (68.99)	-89.18 (82.67)	6.20 (9.49)	-3.28 (2.17)
Avg. Control	149.93 (52.03)	360.11 (48.92)	25.32 (3.96)	9.80 (1.64)
Observations	240	240	224	240
Milei and Bullrich only	-45.82 (62.70)	-125.78 (98.32)	-12.81 (11.56)	-4.35** (1.76)
Avg. Control	159.83 (45.22)	328.63 (73.17)	35.46 (9.95)	8.98 (1.27)
Observations	429	429	390	429
Least - most favourite	-59.34 (67.65)	-131.30 (103.48)	-2.83 (12.57)	-5.91*** (1.97)
Avg. Control	205.50 (46.13)	545.10 (71.08)	53.83 (9.22)	14.63 (1.42)
Observations	515	515	470	515

Notes: The table reports coefficients from regressions of forecast-gap measures on the information-treatment indicator, restricted to respondents who did not receive the accuracy treatment. The main coefficients are in the first row, where the outcome is the maximum difference across the three candidate-contingent forecasts for each indicator. The remaining rows report pairwise candidate gaps and the least-minus-most preferred candidate gap. Robust standard errors are in parentheses. The control-group mean is reported for each outcome. *, **, and *** indicate significance at the 10%, 5%, and 1% levels.

but the effects are generally weaker than in the accuracy-treatment condition. The most consistent reductions again appear for unemployment. The unemployment gap falls significantly in the Massa–Milei and Milei–Bullrich comparisons, and the least-versus-most preferred candidate gap declines by 5.9 percentage points. By contrast, the estimated effects for the exchange rates and inflation are less stable and generally imprecisely estimated.

The CDFs in Figure 3 provide distributional evidence consistent with this interpretation. In the comparison between the information treatment and the control group among respondents without accuracy incentives, the clearest distributional shift is for unemployment. The Kolmogorov–Smirnov test rejects equality of distributions for unemployment, while the differences for the official dollar, blue dollar, and inflation are not statistically significant (see Table OA.9). Thus, information provision appears to move forecast gaps in the expected direction, but the distributional evidence is concentrated in the outcome for which respondents were least informed about current levels.

Figure 3 : Distribution of Maximum Forecast Gaps by Information-Treatment Status



Notes: The figure plots empirical cumulative distribution functions of the maximum forecast gap for each economic indicator, separately by information-treatment status. The sample is restricted to respondents who did not receive the accuracy treatment. The maximum forecast gap is defined as the difference between the highest and lowest candidate-contingent forecast across the three candidates. Corresponding Kolmogorov–Smirnov tests are reported in Table OA.9.

Table 4 examines whether the information treatment reduces dispersion in the forecast

levels themselves.¹⁹ The evidence is mixed but informative. For the Massa scenario, the standard deviation falls substantially for the official dollar and inflation. For the Bullrich scenario, dispersion falls for the official dollar, the blue dollar, and inflation. These reductions are statistically significant and economically meaningful. For the Milei scenario, however, there is no statistically significant reduction in dispersion; if anything, the standard deviations of the exchange-rate forecasts increase. The information treatment also does not significantly reduce dispersion in unemployment forecasts for any candidate scenario.

Table 4—: Information Provision and Forecast Dispersion

Massa					
	Control	Treatment	Difference	P-value	Obs.
Official dollar	675.36	479.03	-196.32***	0.000	531
Dollar blue	814.65	832.15	17.49	0.730	531
Inflation	154.88	119.53	-35.35***	0.000	485
Unemployment	29.36	28.75	-0.60	0.736	531
Milei					
	Control	Treatment	Difference	P-value	Obs.
Official dollar	800.91	881.86	80.95	0.118	531
Dollar blue	1172.88	1223.48	50.60	0.492	531
Inflation	138.75	127.37	-11.38	0.184	485
Unemployment	26.25	26.12	-0.12	0.939	531
Bullrich					
	Control	Treatment	Difference	P-value	Obs.
Official dollar	686.44	512.01	-174.43***	0.000	531
Dollar blue	766.74	619.17	-147.57***	0.001	531
Inflation	133.92	87.35	-46.58***	0.000	485
Unemployment	25.44	24.99	-0.46	0.769	531

Notes: The table compares the dispersion of candidate-contingent forecasts between respondents who received the information treatment and respondents who did not, restricted to respondents who did not receive the accuracy treatment. Each panel corresponds to forecasts under the candidate named in the panel heading. “Control” reports the standard deviation of forecasts in the no-information group, and “Treatment” reports the standard deviation in the information-treatment group. “Difference” is Treatment minus Control, so negative values indicate lower dispersion after information provision. P-values are from two-sample tests of equality of standard deviations. Obs. reports the number of non-missing observations used in each comparison. *, **, and *** indicate significance at the 10%, 5%, and 1% levels.

¹⁹Table OA.8 shows the standard deviations of the forecasts by each type of supporter and candidate.

Overall, information provision seems to have effects that are directionally similar to those of accuracy incentives, but weaker and less systematic. This is consistent with the idea that factual information can partially anchor forecasts, especially when respondents are poorly informed about current conditions, but that information alone is not enough to consistently undo politically motivated differences in expectations. In a highly unstable environment such as Argentina, respondents may already be attentive to salient indicators such as exchange rates (Weber et al., 2025; Cavallo et al., 2017), which may help explain why the information treatment has more limited effects than monetary incentives.

D. Combined treatments

The 2x2 design allows me to examine whether accuracy incentives add to the effect of providing information about current economic conditions. This is relevant because information provision is simpler and cheaper to implement in surveys than monetary incentives and helps distinguish whether accuracy incentives operate partly by correcting misinformation. I estimate specifications that include indicators for the information treatment, the accuracy treatment, and their interaction. Table 5 shows that the two interventions do not appear to be strong complements. The accuracy treatment is the main driver of the reduction in forecast gaps, especially for the blue dollar and unemployment. Among respondents who did not receive information, accuracy incentives reduce the blue-dollar gap by 206 pesos and the unemployment gap by 4.9 percentage points. Information provision also moves the maximum gaps in the expected direction for the blue dollar, inflation, and unemployment, but its effects are smaller and less precise. The only statistically significant effect of information alone is for unemployment, where the gap falls by 3.7 percentage points.

The interaction term is only statistically significant for unemployment. The positive sign indicates that although each treatment reduces the unemployment gap when introduced separately, the combined treatment produces a smaller reduction than the sum of the separate effects, indicating some interference. Relative to the pure control group, receiving both treatments reduces the unemployment gap by 2.8 percentage points,

Table 5—: Factorial Treatment Effects on Economic-Indicator Forecast Gaps

	(1) Official dollar	(2) Dollar blue	(3) Inflation	(4) Unemployment
Accuracy treatment	-26.12 (58.76)	-206.30** (82.47)	-18.23 (12.50)	-4.94*** (1.59)
Information treatment	23.49 (63.72)	-52.83 (101.04)	-6.62 (13.11)	-3.74** (1.78)
Accuracy x information	-64.21 (85.15)	58.88 (117.86)	12.09 (17.21)	5.86** (2.39)
Avg. Control	553.22	871.09	81.08	20.33
Avg. Accuracy only	527.10	664.79	62.85	15.39
Avg. Information only	576.71	818.25	74.46	16.59
Avg. Both treatments	486.38	670.83	68.31	17.51
Observations	1058	1055	951	1058

Notes: The table reports coefficients from regressions of the maximum forecast gap for each economic indicator on the accuracy-treatment indicator, the information-treatment indicator, and their interaction. The maximum forecast gap is defined as the difference between the highest and lowest candidate-contingent forecast across the three candidates. The omitted group is the pure control group, which received neither accuracy incentives nor information about current economic conditions. Average rows report mean maximum gaps by treatment cell. Robust standard errors are in parentheses. *, **, and *** indicate significance at the 10%, 5%, and 1% levels.

compared with reductions of 4.9 percentage points under accuracy incentives alone and 3.7 percentage points under information alone. This suggests that, for an indicator about which respondents are initially less informed, providing current information already captures part of what accuracy incentives would otherwise achieve.

Figure OA.6 provides a distributional comparison between the pure control group and the combined-treatment group; the Kolmogorov–Smirnov tests reject equality of distributions only for unemployment gaps (see Table OA.9), consistent with the regression evidence. Taken together, the 2x2 results suggest that information provision can account for part of the reduction in forecast gaps when respondents are poorly informed about current conditions, but that accuracy incentives are the main driver of the overall reduction in partisan forecast gaps.

E. Heterogeneity

I next explore whether the treatment effects vary with respondents' literacy, prior knowledge about the current economic conditions, attention, and political orientation.²⁰ These analyses are descriptive, since they split the sample into smaller cells, but they are useful for understanding whether the main results are concentrated among particular types of respondents. The figures compare three mutually exclusive groups: the pure control group, the accuracy-treatment-only group, and the information-treatment-only group. Respondents who received both treatments are excluded from these figures so that each treatment arm is compared with the same control group.

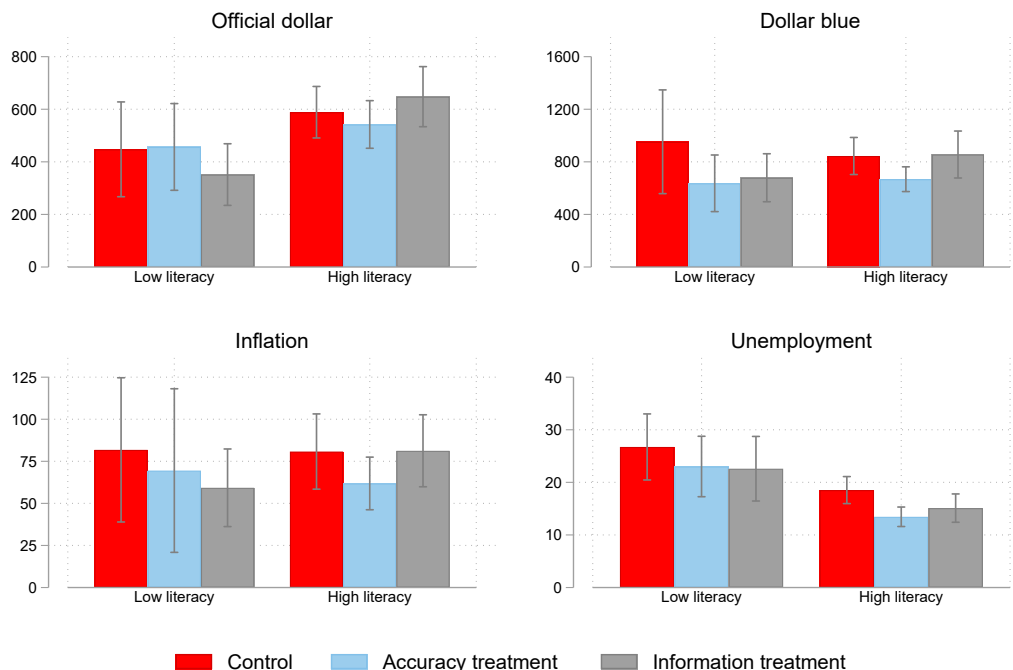
Figure 4 splits respondents by literacy score. Accuracy incentives reduce maximum forecast gaps among both low- and high-literacy respondents, especially for the blue-dollar and unemployment gaps. Information provision appears more concentrated among low-literacy respondents. This is consistent with the interpretation that information provision is most useful when respondents enter the forecasting task with weaker prior knowledge.

Figure 5 splits respondents according to how accurately they reported the current value of each indicator. Accuracy incentives reduce gaps among respondents with both high and low current-value accuracy, especially for the blue dollar and the unemployment gaps. Information provision shows weaker effects, but reductions are generally stronger among respondents with less accurate current beliefs. This suggests that providing correct current values matters most when respondents initially have less accurate benchmarks, but the effect of accuracy incentives does not depend mechanically on respondents' knowledge of current conditions.

The time spent on the forecasting task increased with the accuracy treatment but not with the information treatment (Figure OA.7). Figure 6 examines heterogeneity by the time respondents spent on the candidate-contingent forecast page. Accuracy incentives reduce blue-dollar gaps among both faster and slower respondents. The reduction in inflation and unemployment gaps is more visible among slower respondents. Information

²⁰The Online Appendix Figures OA.8 and OA.9 explore heterogeneity by age, gender. These splits were not pre-registered.

Figure 4 : Maximum Forecast Gaps by Literacy



Notes: The figure reports mean maximum forecast gaps by treatment arm and literacy group. The maximum forecast gap is defined as the difference between the highest and lowest candidate-contingent forecast across the three candidates. The three treatment arms are mutually exclusive: pure control, accuracy treatment only, and information treatment only. Respondents who received both treatments are excluded. Low-literacy respondents answered 0–4 literacy questions correctly, and high-literacy respondents answered 5–9 correctly. Vertical bars show 95% confidence intervals.

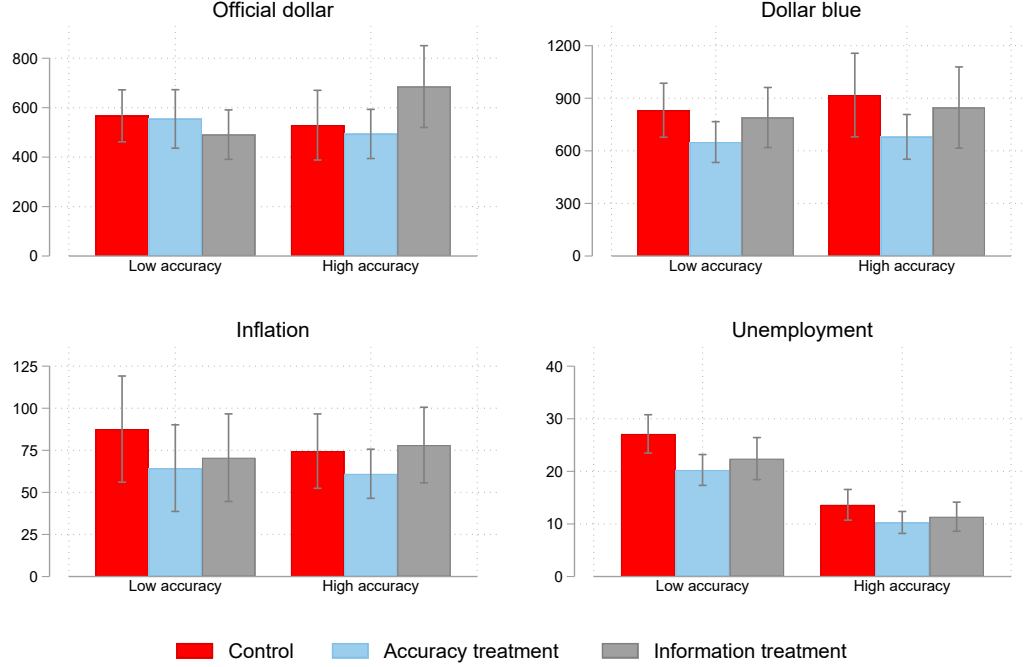
provision also appears to reduce gaps more clearly among slower respondents, especially for unemployment.²¹

Finally, Figure 7 separates respondents by political leaning. Accuracy incentives reduce forecast gaps among both left- and right-leaning respondents, with particularly large reductions among left-leaning respondents for the blue dollar and unemployment. Information provision tends to reduce gaps more clearly among right-leaning respondents.

The candidate-support figures in the Online Appendix (Figures OA.11, OA.13 and OA.12) show the same broad pattern within supporter groups: respondents report more

²¹The current-values page took a median of 174 seconds. Figure OA.10 shows that respondents who spent more time on this page reported higher confidence and answered more current-value questions correctly.

Figure 5 : Maximum Forecast Gaps by Current-Value Accuracy



Notes: The figure reports mean maximum forecast gaps by treatment arm and accuracy regarding current values of the indicators. The maximum forecast gap is defined as the difference between the highest and lowest candidate-contingent forecast across the three candidates. For each indicator, current-value accuracy is measured as the absolute distance between the respondent's reported current value and the benchmark value. High-accuracy respondents are those below the median distance for that indicator, and low-accuracy respondents are those at or above the median. Respondents who received both treatments are excluded. Vertical bars show 95% confidence intervals.

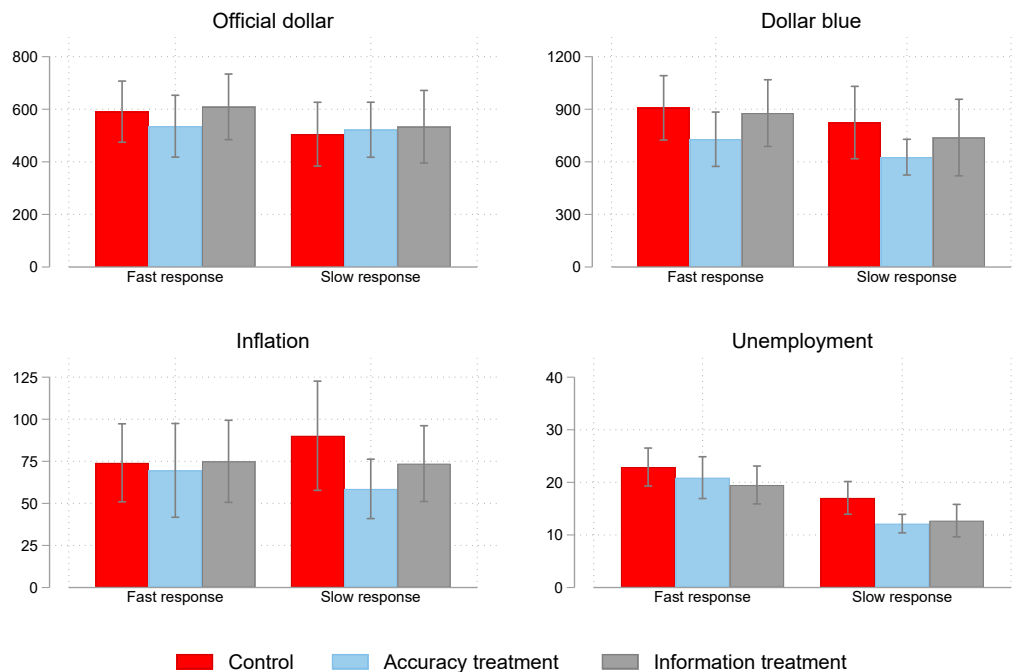
favorable forecasts under their preferred candidate, and the accuracy and information treatments tend to compress these within-supporter differences.

Overall, the heterogeneity analysis suggests that the main findings are not driven by a single subgroup. Accuracy incentives reduce partisan forecast gaps among respondents with different levels of literacy, current-value knowledge, response times, and political orientations. Information provision has weaker and less systematic effects, but appears more useful among respondents with less accurate prior information.

F. Robustness

I next examine whether the results are sensitive to extreme responses and attention.

Figure 6 : Maximum Forecast Gaps by Response Time

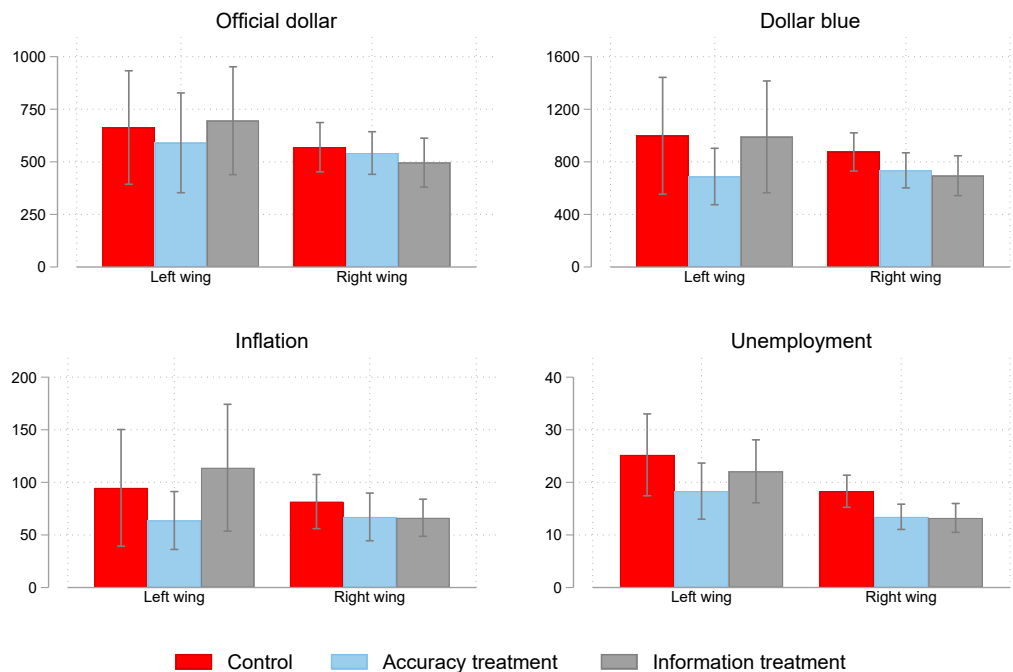


Notes: The figure reports mean maximum forecast gaps by treatment arm and response-time group. The maximum forecast gap is defined as the difference between the highest and lowest candidate-contingent forecast across the three candidates. Fast-response respondents are those who spent at or below the median time on the candidate-contingent forecast page, and slow-response respondents are those who spent above the median time. Respondents who received both treatments are excluded. Vertical bars show 95% confidence intervals.

The first concern is especially relevant for inflation and the exchange-rate forecasts, which were not naturally bounded, unlike unemployment. I therefore re-estimate the main treatment-effect specifications after winsorizing the underlying candidate-level forecasts at the 5th and 95th percentiles. The gap variables are then recomputed using these winsorized forecasts. I do this separately for the accuracy treatment, restricting the comparison to respondents who did not receive the information treatment, and for the information treatment, restricting the comparison to respondents who did not receive the accuracy treatment.

The results for the accuracy treatment are broadly robust (Table OA.16). After winsorization, the treatment still reduces the maximum forecast difference for the blue dollar and unemployment. The estimated effect is -102.86 for the blue-dollar gap and

Figure 7 : Maximum Forecast Gaps by Political Leaning



Notes: The figure reports mean maximum forecast gaps by treatment arm and political-leaning group. The maximum forecast gap is defined as the difference between the highest and lowest candidate-contingent forecast across the three candidates. Left-leaning respondents are those with scores from 0 to 4 on the right-wing policy preference scale, and right-leaning respondents are those with scores from 6 to 10. Respondents with a score of 5 and respondents who received both treatments are excluded. Vertical bars show 95% confidence intervals.

−4.40 for the unemployment gap, both statistically significant. The least-minus-most-favourite gap also falls for all four outcomes, with significant effects for the blue dollar and unemployment.

The information-treatment results also show that the main findings are not driven by outliers (Table OA.17). Winsorizing the forecasts leaves a strong reduction in the unemployment gaps: the maximum difference falls by 3.69, and the least-minus-most-favourite gap falls by 5.18, both statistically significant. The effect on the maximum blue-dollar difference is also negative, although not always statistically significant. As with the accuracy treatment, there is little evidence of robust effects for inflation, and the official-dollar results are generally imprecise.

Overall, winsorizing the forecasts attenuates some estimates but does not change

the main interpretation. The strongest and most robust effects are concentrated in the blue-dollar and unemployment forecasts, while the inflation results remain weak. This suggests that the main findings are not an artifact of a small number of extreme forecasts.

The other concern with online surveys is inattention. The survey had two attention checks at the very end of the survey. Because a long questionnaire followed the experiment, these failures should be taken as a rather strict test. Appendix Table OA.18 shows that attention-check failures (around 12.5%) are not concentrated in any particular treatment arm, a chi-squared test does not reject equality across the four cells ($p = 0.494$). The corresponding rates for failing both attention checks are also similar across treatment arms, ranging from 4.2% to 7.8% ($p = 0.381$). As a robustness check, Appendix Tables OA.19 and OA.20 re-estimate the main treatment-effect specifications after restricting the sample to respondents who passed both attention checks. The qualitative patterns are broadly similar and the restricted-sample estimates do not reveal a different pattern of treatment effects. However, the restriction reduces the sample to 733 respondents and some estimates, especially in the accuracy treatment, become smaller and less precisely estimated. We therefore interpret the attention-check robustness exercise as evidence that the main conclusions are not driven by differential attention-check failure across treatment cells, while noting that the restricted-sample estimates are noisier.

G. Placebos

In order to check whether there are any mechanical reasons driving the results when we give incentives for accuracy, I asked the subjects to also make forecasts about other variables, namely, whether the MERVAL index and the price of the Bitcoin will go up or down on April 1st, the average temperature in Corrientes on March 1st and what will the football teams with the most and least points in the national league by April 1st.²² A priori, we would expect to only find a significant effect for the football question,

²²Participants who received accuracy incentives for the economic forecasts did not receive accuracy incentives for these placebo forecasts, and participants who did not receive accuracy incentives for the economic forecasts did receive incentives for the placebo forecasts. See Figure 1.

but not on the other three.²³ These questions were not candidate-contingent. Table OA.21 presents the results of the regressions of each of these forecasts on the accuracy treatment.

The forecasts for the Merval index and Bitcoin price behave as expected, i.e., no significant differences between treatment and control. Surprisingly, the effects of football were also non-significant. A potential explanation for this is that some respondents do not follow football, so the bias may be weak (even if they have a favorite team), and/or other people who follow football but are also very knowledgeable about the teams' possibilities and might also face a cognitive dissonance effect. This result illustrates the strength of the partisan bias in political topics as compared to other highly passionate contexts such as football in Argentina.

The forecasts regarding the temperature in Corrientes present a statistically significant effect of -1.6°C (equivalent to a change from 90°F to 87°F), which is closer to the actual value of 26.38°C . This result is unexpected, but it is consistent with incentives inducing respondents to think more carefully about a question that may otherwise elicit intuitive or exaggerated answers, especially in a region associated with high temperatures. Since the temperature question was non-contingent and has no obvious partisan content, this effect does not appear to reflect partisan bias. Rather, it suggests that accuracy incentives can improve some simple factual forecasts by encouraging more deliberation.

VI. Discussion

This paper studies whether political preferences shape economic forecasts. The main result is that respondents report more favorable economic outcomes under their preferred candidate, and that monetary incentives for accuracy reduce these candidate-contingent forecast gaps. The effects are strongest for the blue dollar and unemployment. Information about current economic conditions also moves some forecasts in the expected direction and reduces dispersion, but it has weaker and less systematic effects

²³For the football question, I only use the best team as there is almost no variation in the worst team and a significant number of missing values.

on partisan gaps in the forecasts.

A central issue is whether these results reflect a reduction in partisan bias or simply reduced noise in forecasts. The evidence does not support a pure noise interpretation. The design asks the same respondent to forecast the same indicators under different candidate scenarios, which absorbs many differences in economic knowledge, optimism, and attention. Moreover, the gaps are politically directional: respondents tend to forecast better outcomes under their preferred candidate and worse outcomes under candidates they do not support. Random imprecision should not systematically favor the respondent's own candidate. The pattern is similar in the responses to the current values, where the information about the incumbent candidate is available. This pattern is also visible in the pairwise comparisons and in the least-versus-most preferred candidate gaps. The decomposition exercise further shows that accuracy incentives mainly reduce pessimism about non-preferred candidates, rather than simply moving all forecasts toward a common benchmark.

The comparison between the accuracy and information treatments is also informative. If the forecast gaps were driven only by lack of knowledge about the economy, then providing correct current values should substantially reduce them. The evidence is not consistent with this interpretation alone. Information provision reduces forecast dispersion in some cases, suggesting that inaccurate or heterogeneous starting points matter for precision. However, it does not systematically remove the politically aligned gaps between candidate-contingent forecasts.

The blue-dollar results are especially informative. Respondents are highly knowledgeable about the current blue-dollar value: their reported current values are close to the benchmark, as this is the most salient indicator in everyday economic discussions in Argentina. Yet the accuracy treatment produces one of its largest effects precisely for this outcome. This makes it unlikely that the blue-dollar gap is driven mainly by ignorance about economic conditions. Instead, the result suggests that even when respondents share a relatively accurate starting point, political preferences can still shape how they forecast the future.

These findings should not be interpreted as showing that incentives reveal respon-

dents’ latent true beliefs. Rather, they show that unincentivized forecasts contain a politically directional component that becomes smaller when accuracy is rewarded. This reduction may reflect reduced expressive responding, discourage extreme answers, increase effort or induce information search (Graham, 2025), or make respondents rely less on partisan heuristics. The interpretation therefore depends on whether incentives mainly change what respondents report or what they come to believe.

This distinction is related to the debate over motivated beliefs versus “cheerleading” in survey responses (Bullock and Lenz, 2019; Peterson and Iyengar, 2021). If respondents privately hold less partisan forecasts but report more partisan answers when forecasts are unincentivized, the results reflect cheerleading or expressive responding (Graham, 2025; Mian et al., 2023; Bullock et al., 2015). If, instead, accuracy incentives lead respondents to reassess the likelihood of different economic scenarios, the results would reflect motivated beliefs rather than pure expressive responding (e.g., Mayraz, 2011).

The present design cannot fully distinguish between these channels. For the measurement of expectations, however, the central implication is the same: in unincentivized surveys, elicited forecasts contain a politically directional component that becomes smaller when accuracy is rewarded. The evidence is consistent with both mechanisms. Relatively modest incentives reduce forecast gaps, and the adjustment is concentrated in forecasts for non-preferred candidates, which is consistent with an expressive component. At the same time, the treatment-intensity results show that respondents respond to the perceived stakes of accuracy, and the persistence of gaps under incentives suggests that some differences may reflect genuine beliefs about candidates’ economic consequences.

For the interpretation of expectation surveys, both mechanisms are consequential. If unincentivized responses partly reflect cheerleading, accuracy incentives help recover forecasts that are more decision-relevant. If partisan gaps instead arise from motivated beliefs, the effect of incentives remains informative because it shows how reported expectations adjust when accuracy becomes salient. This matters beyond survey measurement: prior evidence links partisan economic expectations to consumption behavior (Gerber and Huber, 2009), and recent work links partisan alignment to high-stakes

economic behavior, including home purchases (Cao and Wu, 2023), entrepreneurship (Engelberg et al., 2026), and professional financial judgments (Kempf and Tsoutsoura, 2021).

The external validity of the results depends on both macroeconomic salience and political salience. Recent evidence shows that attention to inflation is endogenous to the economic environment: consumers pay more attention when inflation is high (Bracha and Tang, 2025), and households and firms in high-inflation settings respond less to inflation-information treatments than agents in low-inflation settings (Weber et al., 2025). My results point to a similar logic within a single high-inflation setting (Argentina in 2023). Respondents were not equally informed about all macroeconomic indicators. They were highly accurate about the current blue dollar, an indicator reported daily and central to everyday economic discussion, but much less accurate about unemployment. Thus, salience varies across indicators even when the macroeconomic environment is unstable. This helps interpret why information provision has limited effects for some outcomes and stronger effects for unemployment, while accuracy incentives reduce gaps even for the blue dollar. The latter result is particularly informative: partisan bias appears not only when respondents lack basic information, but also when they start from a relatively accurate and commonly known current value.

The magnitudes should therefore not be extrapolated mechanically to other countries or periods. External validity depends on both macroeconomic and political salience. On the macroeconomic side, the cross-country evidence in Weber et al. (2025) shows that households and firms in high-inflation environments are more attentive to inflation and respond less to information treatments than those in low-inflation environments. This suggests that information provision may have larger effects in more stable economies, where macroeconomic variables are less salient, while in high-inflation settings its effects may be concentrated on less visible indicators such as unemployment. On the political side, the partisan component should be stronger when elections are polarized, parties propose sharply different economic programs, and voters view economic outcomes as signals of political competence (Gerber and Huber, 2010). Consistent with this, evidence from the United States shows that partisan identity affects inflation and economic

expectations even in a more stable macroeconomic environment ([Coibion et al., 2020](#); [Mian et al., 2023](#); [Binder et al., 2024](#); [DiGiuseppe et al., 2025](#)).

The results have implications for the measurement of economic expectations. Household expectation surveys are often used as inputs for policy analysis, but the findings here suggest that political preferences can affect elicited forecasts when outcomes are politically salient. This matters especially in polarized environments, where aggregate expectations may partly reflect the political composition of the sample. The experiment shows that small monetary incentives can significantly improve the elicitation of expectations. Future work could experimentally vary the size of incentives, measure information acquisition during the forecasting task, and test whether similar patterns arise in representative samples and in more stable macroeconomic environments.

REFERENCES

- Allcott, Hunt and Matthew Gentzkow**, “Social media and fake news in the 2016 election,” *Journal of Economic Perspectives*, 2017, 31 (2), 211–236.
- Armantier, Olivier, Scott Nelson, Giorgio Topa, Wilbert Van der Klaauw, and Basit Zafar**, “The price is right: Updating inflation expectations in a randomized price information experiment,” *Review of Economics and Statistics*, 2016, 98 (3), 503–523.
- , **Wändi Bruine de Bruin, Simon Potter, Giorgio Topa, Wilbert Van Der Klaauw, and Basit Zafar**, “Measuring inflation expectations,” *Annual Review of Economics*, 2013, 5 (1), 273–301.
- Bachmann, Oliver, Klaus Gründler, Niklas Potrafke, and Ruben Seiberlich**, “Partisan bias in inflation expectations,” *Public Choice*, 2021, 186 (3), 513–536.
- Bauer, Kevin, Yan Chen, Florian Hett, and Michael Kosfeld**, “Group identity and belief formation: a decomposition of political polarization,” 2023.
- Bénabou, Roland and Jean Tirole**, “Mindful economics: The production, consumption, and value of beliefs,” *Journal of Economic Perspectives*, 2016, 30 (3), 141–164.
- Binder, Carola Conces, Rupal Kamdar, and Jane M Ryngaert**, “Partisan expectations and COVID-era inflation,” *Journal of Monetary Economics*, 2024, 148, 103649.
- Bracha, Anat and Jenny Tang**, “Inflation levels and (in) attention,” *Review of Economic Studies*, 2025, 92 (3), 1564–1594.
- Bullock, John G, Alan S Gerber, Seth J Hill, and Gregory A Huber**, “Partisan bias in factual beliefs about politics,” *Quarterly Journal of Political Science*, 2015, 10 (4), 519–578.

- **and Gabriel Lenz**, “Partisan bias in surveys,” *Annual Review of Political Science*, 2019, *22*, 325–342.
- Canen, Nathan, Chad Kendall, and Francesco Trebbi**, “Unbundling polarization,” *Econometrica*, 2020, *88* (3), 1197–1233.
- Cao, Xinyu and Tianping Wu**, “Partisan Expectations and Behavior in the US Housing Market,” *Available at SSRN 5027026*, 2023.
- Cavallo, Alberto, Guillermo Cruces, and Ricardo Perez-Truglia**, “Inflation expectations, learning, and supermarket prices: Evidence from survey experiments,” *American Economic Journal: Macroeconomics*, 2017, *9* (3), 1–35.
- Coibion, Olivier and Yuriy Gorodnichenko**, “Schumpeter Lecture 2025: The New Causal Macroeconomics of Surveys and Experiments,” *Journal of the European Economic Association*, 2026, *24* (1), 58–81.
- , – , **and Michael Weber**, “Political polarization and expected economic outcomes,” Technical Report, National Bureau of Economic Research 2020.
- , – , **and –**, “Monetary policy communications and their effects on household inflation expectations,” *Journal of Political Economy*, 2022, *130* (6), 1537–1584.
- , – , **and Rupal Kamdar**, “The formation of expectations, inflation, and the phillips curve,” *Journal of Economic Literature*, 2018, *56* (4), 1447–1491.
- D’Acunto, Francesco, Daniel Hoang, Maritta Paloviita, and Michael Weber**, “Effective policy communication: Targets versus instruments,” *Chicago Booth Research Paper*, 2020, (20-38).
- Dietrich, Alexander, Edward S Knotek II, Kristian O Myrseth, Robert W Rich, Raphael Schoenle, and Michael Weber**, “Greater Than the Sum of Its Parts: Aggregate vs. Aggregated Inflation Expectations,” Technical Report, National Bureau of Economic Research 2023.

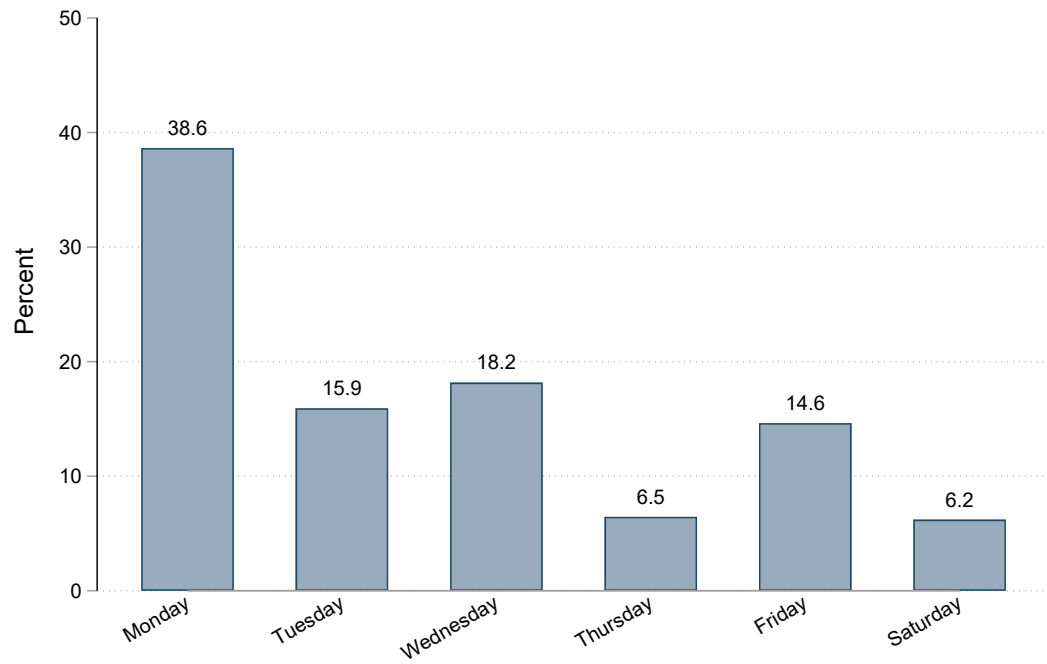
- DiGiuseppe, Matthew, Ana Carolina Garriga, and Andreas Kern**, “Partisan Bias in Inflation Expectations,” MPRA Paper 124391, Munich Personal RePEc Archive 2025.
- Donkor, Kwabena, Lorenz Goette, Maximilian W Müller, Eugen Dimant, and Michael Kurschilgen**, “Identity and Economic Incentives,” 2023.
- D’Acunto, Francesco and Michael Weber**, “Why Survey-Based Subjective Expectations are Meaningful and Important,” Technical Report, National Bureau of Economic Research 2024.
- Ehrbeck, Tilman and Robert Waldmann**, “Why are professional forecasters biased? Agency versus behavioral explanations,” *The Quarterly Journal of Economics*, 1996, *111* (1), 21–40.
- Engelberg, Joseph, Jorge Guzman, Runjing Lu, and William Mullins**, “Partisan Entrepreneurship,” *The Journal of Finance*, May 2026.
- Gemmi, Luca and Rosen Valchev**, “Biased surveys,” *Journal of Monetary Economics*, 2025, p. 103868.
- Gerber, Alan S and Gregory A Huber**, “Partisanship and economic behavior: Do partisan differences in economic forecasts predict real economic behavior?,” *American Political Science Review*, 2009, *103* (3), 407–426.
- and —, “Partisanship, political control, and economic assessments,” *American Journal of Political Science*, 2010, *54* (1), 153–173.
- Graham, Matthew H**, “Expressive Responding and the Economy: The Case of Trump’s Return to Office,” *Journal of Experimental Political Science*, 2025, pp. 1–14.
- Jerit, Jennifer and Jason Barabas**, “Partisan perceptual bias and the information environment,” *The Journal of Politics*, 2012, *74* (3), 672–684.
- Kempf, Elisabeth and Margarita Tsoutsoura**, “Partisan professionals: Evidence from credit rating analysts,” *The Journal of Finance*, 2021, *76* (6), 2805–2856.

- Kunda, Ziva**, “The case for motivated reasoning,” *Psychological Bulletin*, 1990, *108* (3), 480.
- Laster, David, Paul Bennett, and In Sun Geoum**, “Rational bias in macroeconomic forecasts,” *The Quarterly Journal of Economics*, 1999, *114* (1), 293–318.
- Malmendier, Ulrike and Stefan Nagel**, “Learning from inflation experiences,” *The Quarterly Journal of Economics*, 2016, *131* (1), 53–87.
- Mayraz, Guy**, “Wishful Thinking,” 2011. Unpublished manuscript.
- Mian, Atif, Amir Sufi, and Nasim Khoshkhoh**, “Partisan bias, economic expectations, and household spending,” *Review of Economics and Statistics*, 2023, *105* (3), 493–510.
- Mokhtarzadeh, Fatemeh and Luba Petersen**, “Coordinating expectations through central bank projections,” *Experimental Economics*, 2021, *24* (3), 883–918.
- Ottaviani, Marco and Peter Norman Sørensen**, “Reputational cheap talk,” *The RAND Journal of Economics*, 2006, *37* (1), 155–175.
- and —, “The strategy of professional forecasting,” *Journal of Financial Economics*, 2006, *81* (2), 441–466.
- Peterson, Erik and Shanto Iyengar**, “Partisan gaps in political information and information-seeking behavior: motivated reasoning or cheerleading?,” *American Journal of Political Science*, 2021, *65* (1), 133–147.
- Prior, Markus, Gaurav Sood, Kabir Khanna et al.**, “You cannot be serious: The impact of accuracy incentives on partisan bias in reports of economic perceptions,” *Quarterly Journal of Political Science*, 2015, *10* (4), 489–518.
- Rathje, Steve, Jon Roozenbeek, Jay J Van Bavel, and Sander van der Linden**, “Accuracy and social motivations shape judgements of (mis) information,” *Nature Human Behaviour*, 2023, pp. 1–12.

- Robertson, Ronald E, Jon Green, Damian J Ruck, Katherine Ognyanova, Christo Wilson, and David Lazer**, “Users choose to engage with more partisan news than they are exposed to on Google Search,” *Nature*, 2023, *618* (7964), 342–348.
- Thaler, Michael**, “The fake news effect: Experimentally identifying motivated reasoning using trust in news,” *American Economic Journal: Microeconomics*, 2024, *16* (2), 1–38.
- Vosoughi, Soroush, Deb Roy, and Sinan Aral**, “The spread of true and false news online,” *Science*, 2018, *359* (6380), 1146–1151.
- Weber, Michael, Bernardo Candia, Hassan Afrouzi, Tiziano Ropele, Rodrigo Lluberas, Serafin Frache, Brent Meyer, Saten Kumar, Yuriy Gorodnichenko, Dimitris Georgarakos et al.**, “Tell Me Something I Don’t Already Know: Learning in Low-and High-Inflation Settings,” *Econometrica*, 2025, *93* (1), 229–264.
- , **Francesco D’Acunto, Yuriy Gorodnichenko, and Olivier Coibion**, “The subjective inflation expectations of households and firms: Measurement, determinants, and implications,” *Journal of Economic Perspectives*, 2022, *36* (3), 157–184.

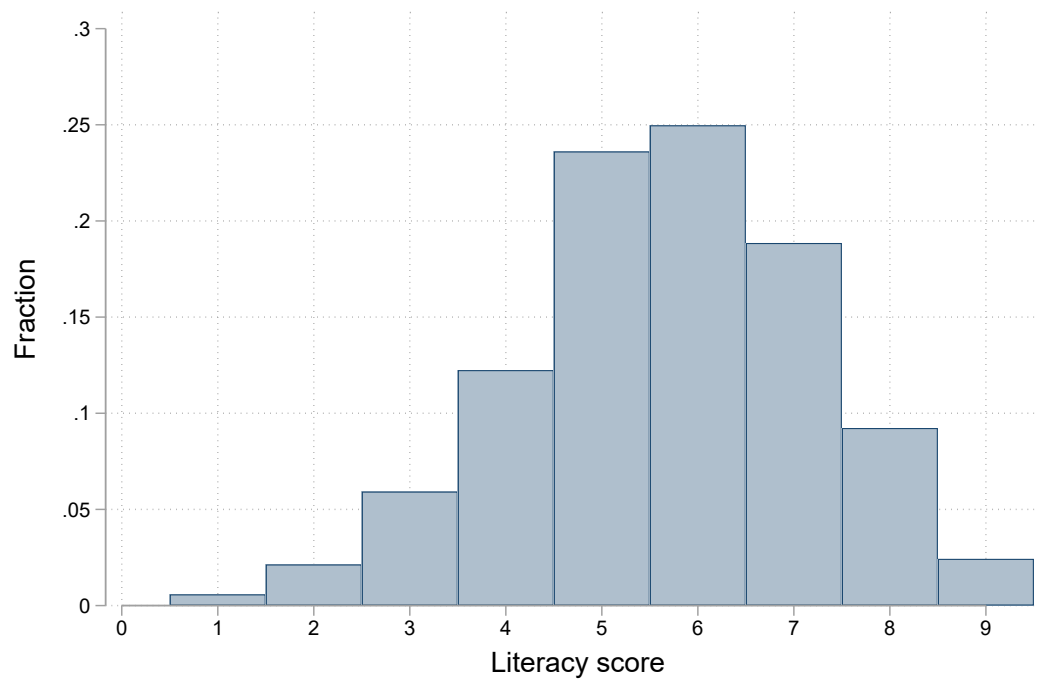
ONLINE APPENDIX**Figures**

Figure OA.1 : Survey Completion Dates



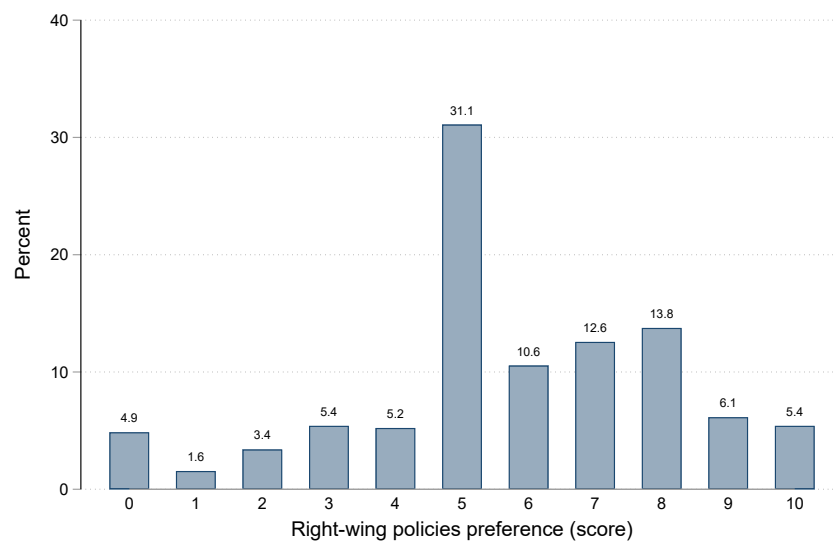
Notes: The figure shows the distribution of survey completion dates during the field period. Bars report the percentage of respondents who completed the survey on each day from Monday, October 16, to Saturday, October 21, 2023.

Figure OA.2 : Literacy Scores



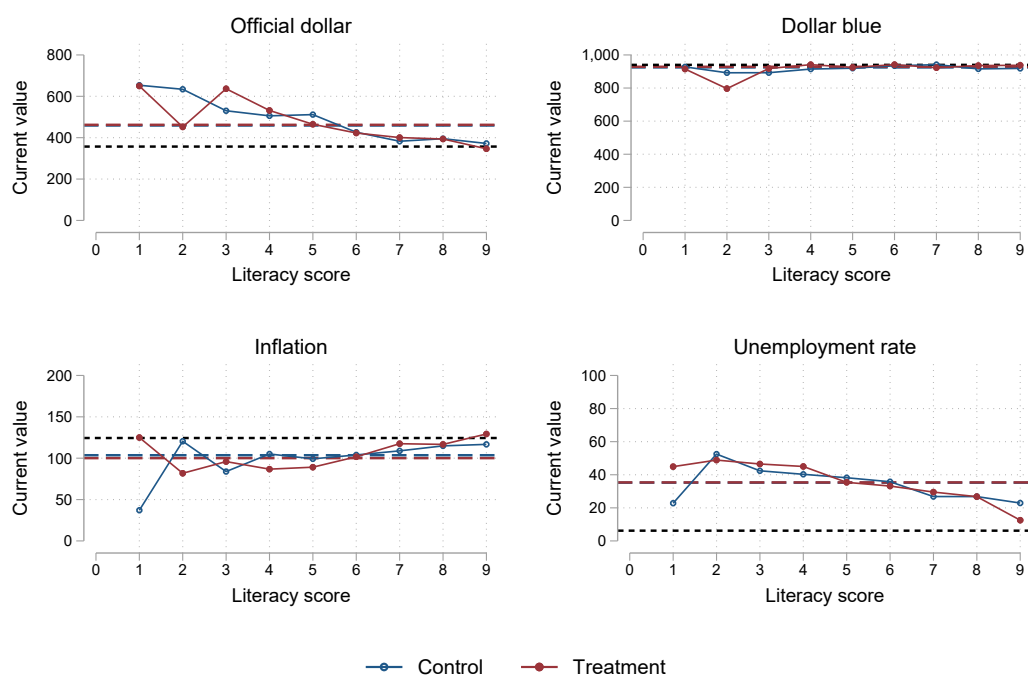
Notes: The figure shows the distribution of literacy scores. The score ranges from 0 to 9 and equals the number of economic and financial literacy questions answered correctly. Bars report the fraction of respondents at each score.

Figure OA.3 : Political Leaning



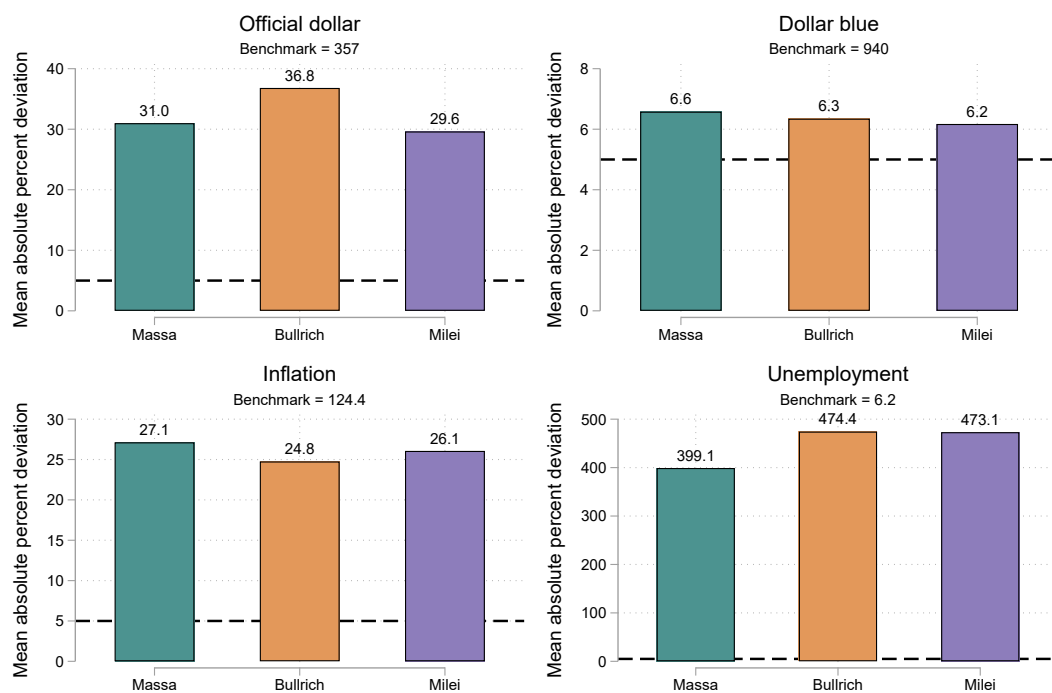
Notes: The figure shows the distribution of responses to the right-wing policy preference question. The score ranges from 0 to 10, with higher values indicating stronger agreement with right-wing policies. Bars report the percentage of respondents at each score.

Figure OA.4 : Reported Current Values by Literacy



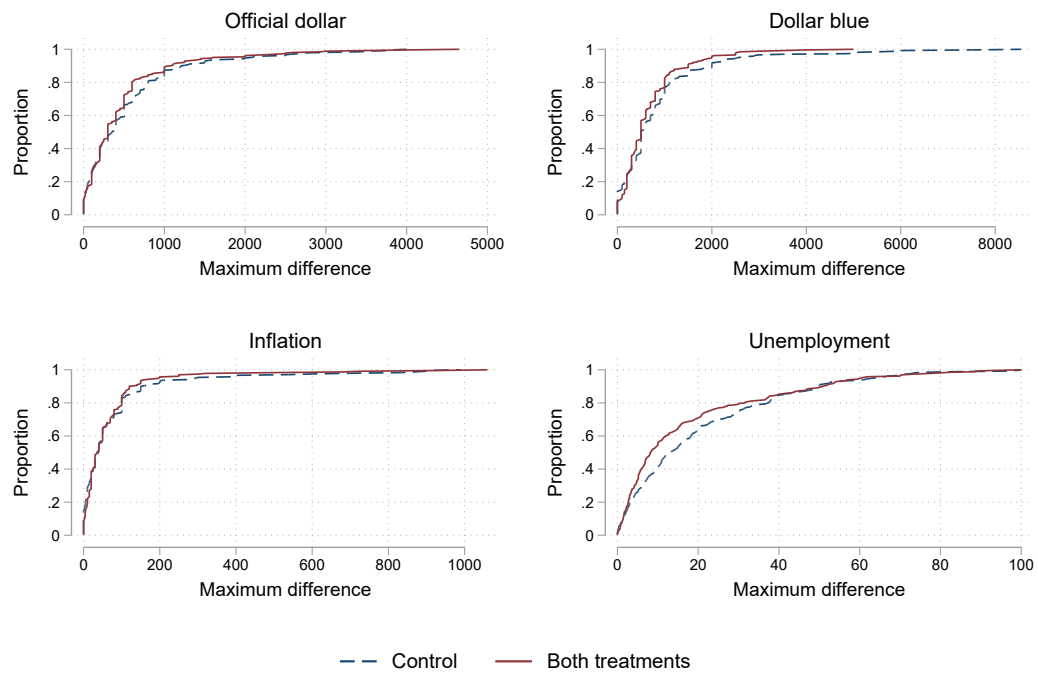
Notes: The figure plots mean reported current values for each economic indicator by literacy score, separately by accuracy-treatment status. The literacy score ranges from 0 to 9 and equals the number of economic and financial literacy questions answered correctly. Dashed horizontal lines report the average response within each treatment group, and solid black horizontal lines report the benchmark values used in the experiment: 357 ARS for the official dollar, 940 ARS for the blue dollar, 124.4% for annual inflation, and 6.2% for unemployment.

Figure OA.5 : Percent Deviations in Current Values by Candidate Support



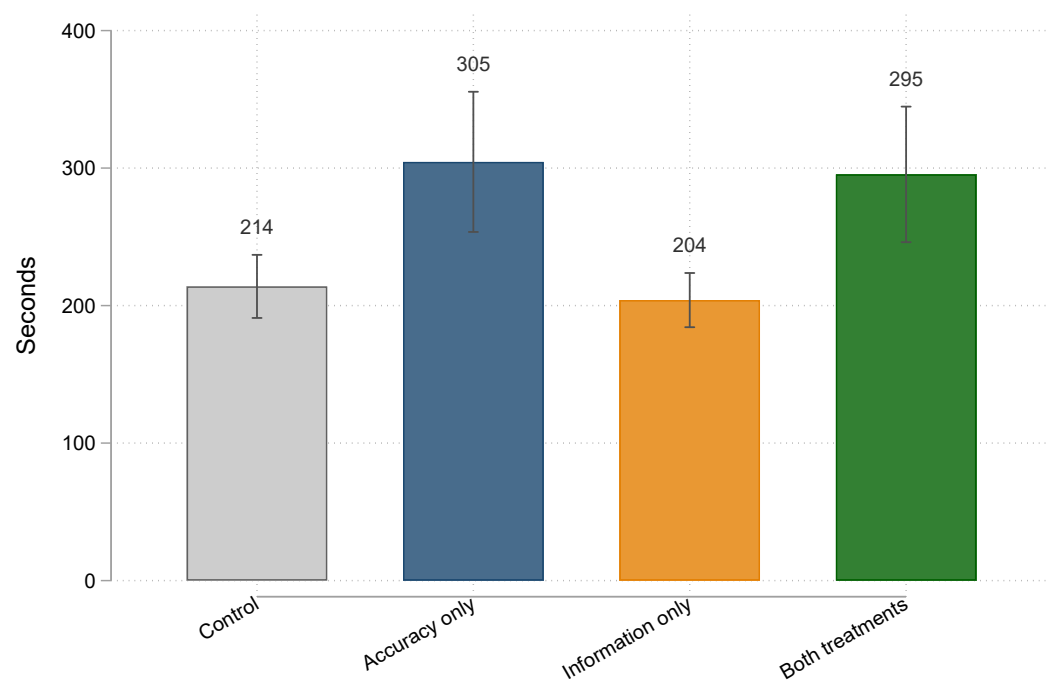
Notes: The figure reports mean absolute percent deviations of respondents' reported current values from the benchmark value for each economic indicator, separately by favorite-candidate support. Percent deviation is defined as $100 \times |\text{reported value} - \text{benchmark value}| / \text{benchmark value}$. Benchmarks are 357 ARS for the official dollar, 940 ARS for the blue dollar, 124.4% for annual inflation, and 6.2% for unemployment. The dashed horizontal line marks the 5% threshold used to classify responses as accurate.

Figure OA.6 : Distribution of Maximum Forecast Gaps: Both Treatments versus Control



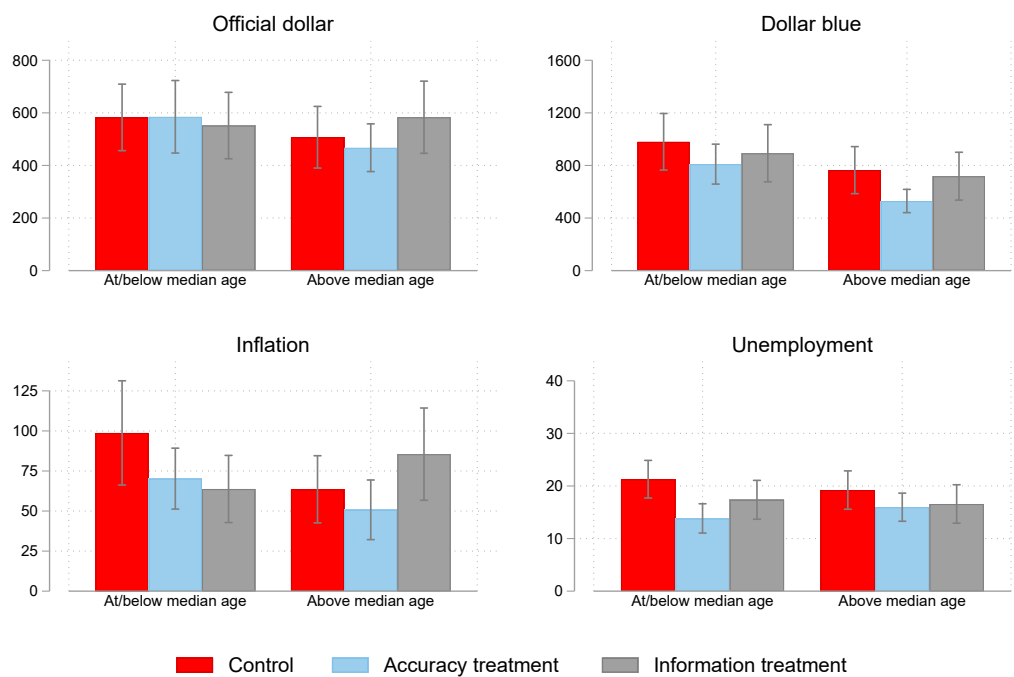
Notes: The figure plots empirical cumulative distribution functions of the maximum forecast gap for each economic indicator, comparing respondents who received both treatments with respondents who received neither treatment. The maximum forecast gap is defined as the difference between the highest and lowest candidate-contingent forecast across the three candidates. Corresponding Kolmogorov–Smirnov tests are reported in Table [OA.9](#).

Figure OA.7 : Forecast Response Time by Treatment Condition



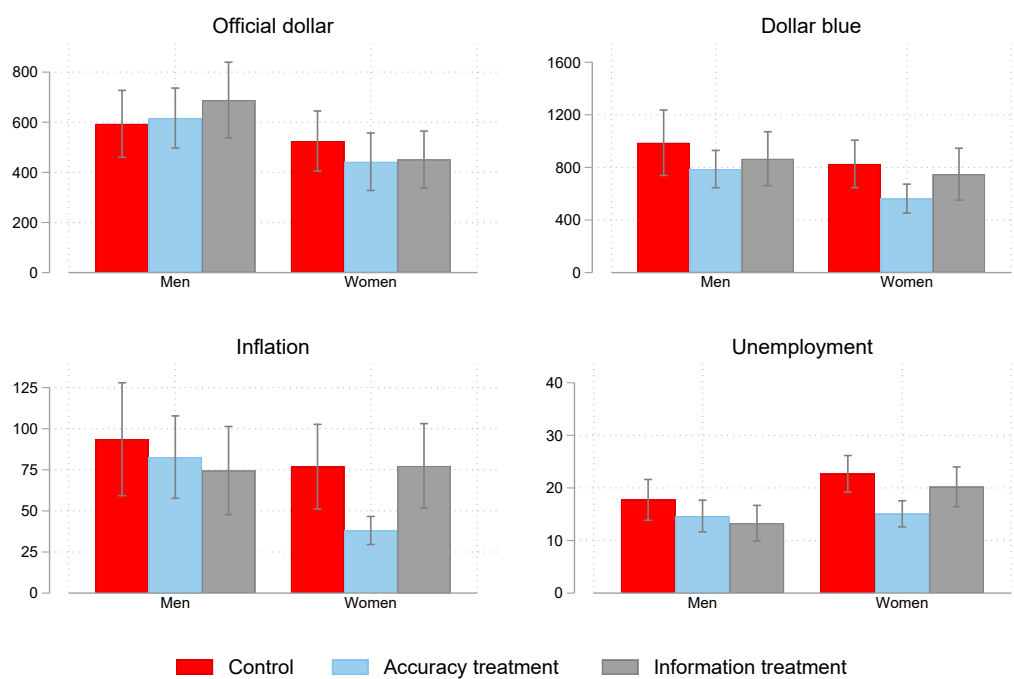
Notes: The figure reports mean time spent on the candidate-contingent forecast page by treatment condition. Time is measured in seconds. Vertical bars show 95% confidence intervals, and labels above the bars report mean response times.

Figure OA.8 : Maximum Forecast Gaps by Age



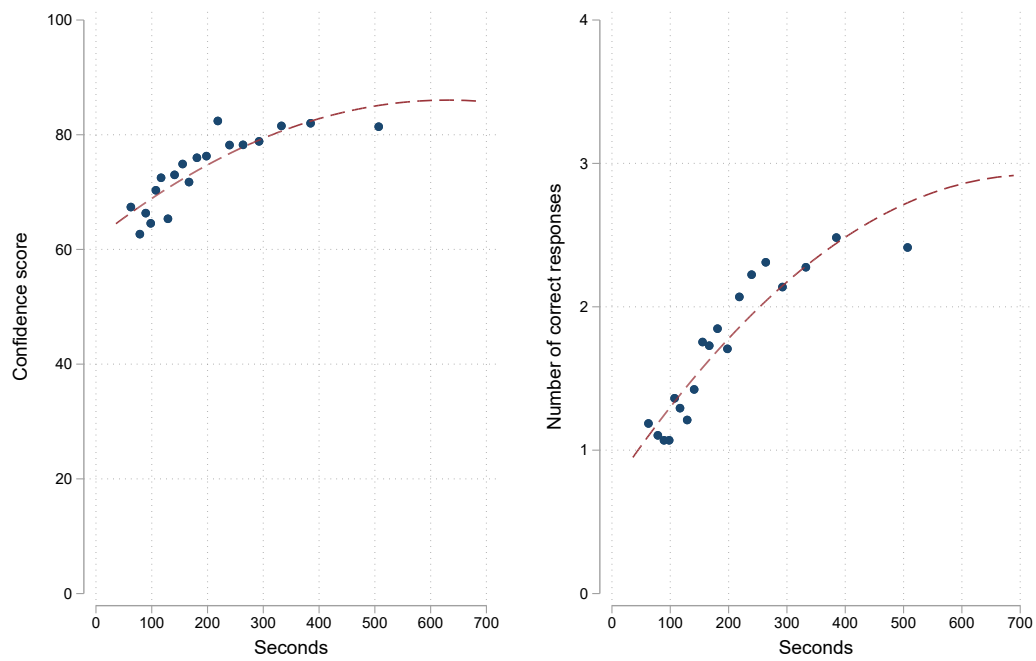
Notes: The figure reports mean maximum forecast gaps by treatment arm and age group. The maximum forecast gap is defined as the difference between the highest and lowest candidate-contingent forecast across the three candidates. The three treatment arms are mutually exclusive: pure control, accuracy treatment only, and information treatment only. Respondents who received both treatments are excluded. Younger respondents are those at or below the median age, and older respondents are those above the median age. Vertical bars show 95% confidence intervals.

Figure OA.9 : Maximum Forecast Gaps by Gender



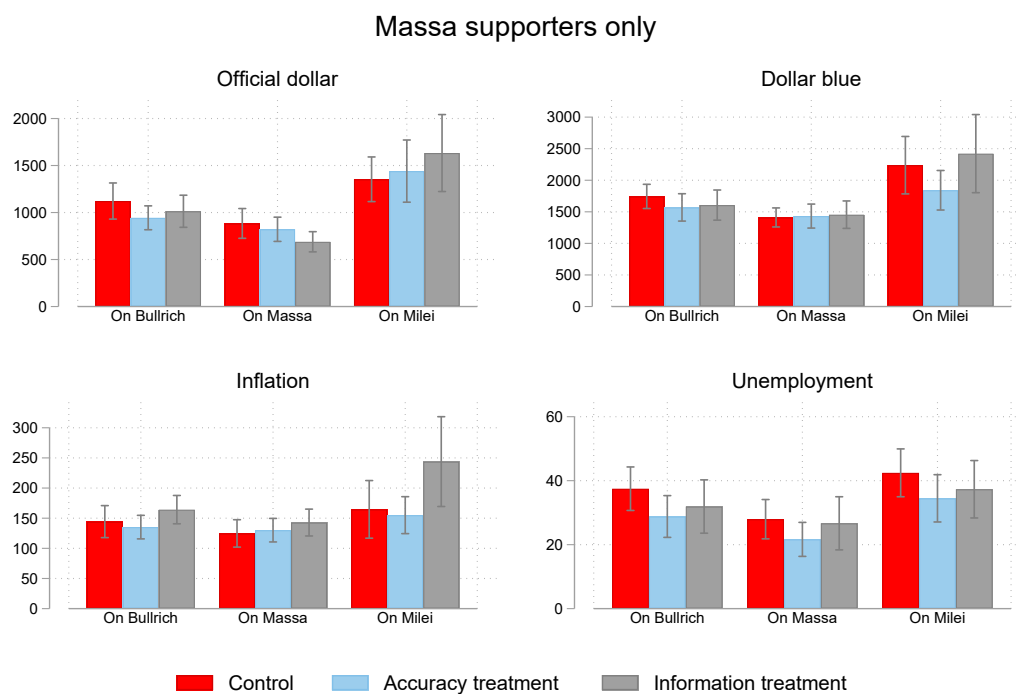
Notes: The figure reports mean maximum forecast gaps by treatment arm and gender. The maximum forecast gap is defined as the difference between the highest and lowest candidate-contingent forecast across the three candidates. The three treatment arms are mutually exclusive: pure control, accuracy treatment only, and information treatment only. Respondents who received both treatments are excluded. The gender split uses the female indicator. Vertical bars show 95% confidence intervals.

Figure OA.10 : Current-Value Confidence and Accuracy by Response Time



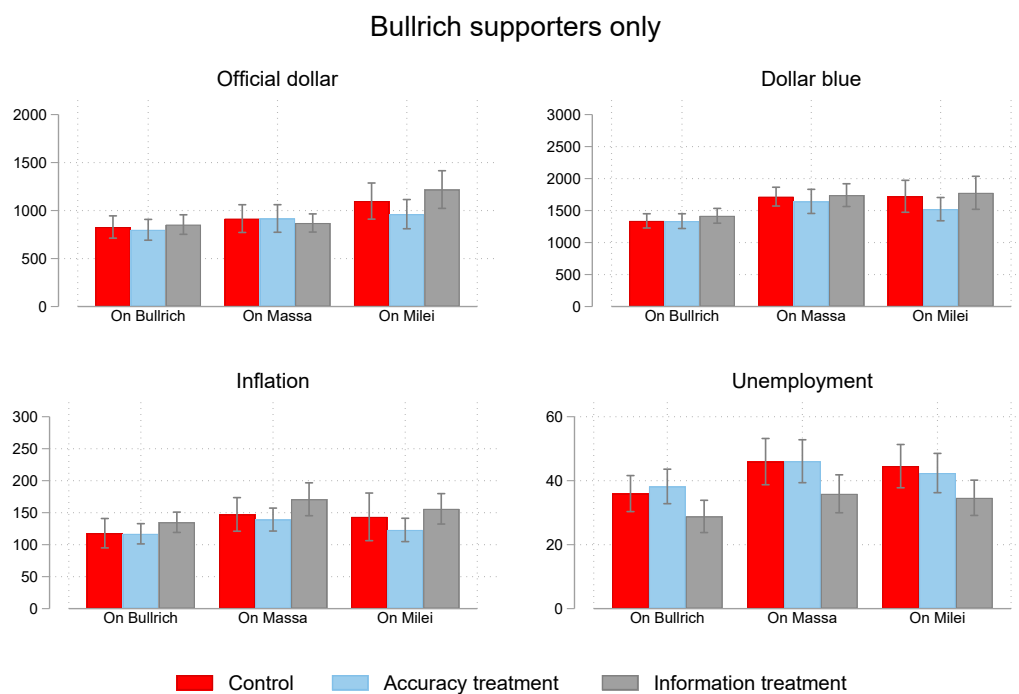
Notes: The figure relates time spent on the current-values page to respondents' confidence and accuracy in reporting current economic conditions. The left panel plots average confidence, measured on a 0–100 scale, against time spent on the page in seconds. The right panel plots the number of current-value questions answered correctly, from 0 to 4, against time spent on the same page. Dots are binned means using 20 bins, and dashed lines show quadratic fits with robust standard errors.

Figure OA.11 : Candidate-Contingent Forecasts among Massa Supporters



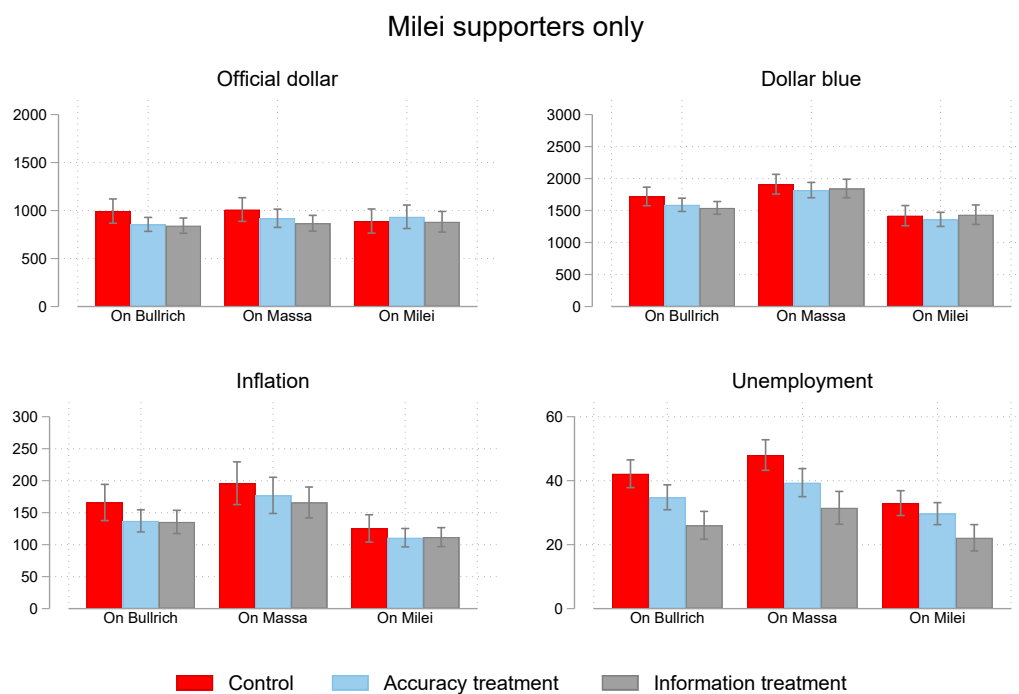
Notes: The figure reports mean candidate-contingent forecasts among Massa supporters only. Each panel corresponds to one economic indicator. Forecasts are shown separately by candidate scenario and treatment arm: pure control, accuracy treatment only, and information treatment only. Respondents who received both treatments are excluded. Exchange-rate forecasts are in ARS, and inflation and unemployment forecasts are in percentage points. Vertical bars show 95% confidence intervals.

Figure OA.12 : Candidate-Contingent Forecasts among Bullrich Supporters



Notes: The figure reports mean candidate-contingent forecasts among Bullrich supporters only. Each panel corresponds to one economic indicator. Forecasts are shown separately by candidate scenario and treatment arm: pure control, accuracy treatment only, and information treatment only. Respondents who received both treatments are excluded. Exchange-rate forecasts are in ARS, and inflation and unemployment forecasts are in percentage points. Vertical bars show 95% confidence intervals.

Figure OA.13 : Candidate-Contingent Forecasts among Milei Supporters



Notes: The figure reports mean candidate-contingent forecasts among Milei supporters only. Each panel corresponds to one economic indicator. Forecasts are shown separately by candidate scenario and treatment arm: pure control, accuracy treatment only, and information treatment only. Respondents who received both treatments are excluded. Exchange-rate forecasts are in ARS, and inflation and unemployment forecasts are in percentage points. Vertical bars show 95% confidence intervals.

Tables

Table OA.1—: Political Preferences by Favorite Candidate

Panel A: Favourite candidate and second preferred candidate				
Favourite candidate	Massa	Bullrich	Milei	N
Massa	0.00	70.24	29.76	168
Bullrich	29.15	0.00	70.85	295
Milei	18.87	81.13	0.00	567
N	193	578	259	1030
Panel B: Favourite candidate and candidate who will vote				
Favourite candidate	Massa	Bullrich	Milei	N
Massa	87.50	4.17	8.33	168
Bullrich	3.05	89.15	7.80	295
Milei	1.23	6.53	92.24	567
N	163	307	560	1030
Panel C: Favourite candidate and candidate who will win				
Favourite candidate	Massa	Bullrich	Milei	N
Massa	37.50	3.57	58.93	168
Bullrich	5.08	25.76	69.15	295
Milei	3.35	3.00	93.65	567
N	97	99	834	1030

Notes: The table reports cross-tabulations between respondents' favorite candidate and other political-preference measures, restricting the sample to respondents who named Massa, Bullrich, or Milei as their favorite candidate. Rows indicate the respondent's favorite candidate. Columns report, respectively, the second preferred candidate, intended vote, and expected winner. Entries are row percentages, and the final column reports the number of respondents in each favorite-candidate group. The bottom row reports column totals.

Table OA.2—: Balance Tests for Demographic Characteristics

	(1) Control		(2) Treatment		(3) Total		(4)
	Mean/SE	N	Mean/SE	N	Mean/SE	N	p-value
Female	0.60 (0.02)	531	0.55 (0.02)	550	0.57 (0.02)	1,081	0.098*
Age	25.95 (0.29)	538	26.05 (0.32)	552	26.00 (0.22)	1,090	0.816
Work	0.57 (0.02)	538	0.59 (0.02)	551	0.58 (0.01)	1,089	0.483
Right-wing policies preference	5.67 (0.11)	538	5.76 (0.10)	552	5.72 (0.07)	1,090	0.536

Notes: The table reports balance tests for demographic characteristics by accuracy-treatment status. Control refers to respondents who did not receive monetary incentives for accurate economic forecasts, and Treatment refers to respondents who did. Columns report means, standard errors in parentheses, and sample sizes. The final column reports p-values from two-sample t-tests of equality of means across treatment groups, allowing unequal variances. *, **, and *** indicate significance at the 10%, 5%, and 1% levels.

Table OA.3—: Balance Tests for Literacy Measures

	(1) Control		(2) Treatment		(3) Total		(4)
	Mean/SE	N	Mean/SE	N	Mean/SE	N	p-value
Monetary literacy score	1.40 (0.03)	514	1.37 (0.03)	515	1.39 (0.02)	1,029	0.373
Financial literacy score	2.01 (0.04)	514	2.02 (0.04)	515	2.02 (0.03)	1,029	0.854
Pure literacy score	2.26 (0.03)	514	2.27 (0.03)	515	2.27 (0.02)	1,029	0.868
Literacy score (all dimensions)	5.68 (0.07)	514	5.66 (0.07)	515	5.67 (0.05)	1,029	0.831

Notes: The table reports balance tests for literacy measures by accuracy-treatment status. Control refers to respondents who did not receive monetary incentives for accurate economic forecasts, and Treatment refers to respondents who did. Columns report means, standard errors in parentheses, and sample sizes. The final column reports p-values from two-sample t-tests of equality of means across treatment groups, allowing unequal variances. *, **, and *** indicate significance at the 10%, 5%, and 1% levels.

Table OA.4—: Balance Tests for Reported Current Values

	(1) Control		(2) Treatment		(3) Total		(4)
	Mean/SE	N	Mean/SE	N	Mean/SE	N	p-value
Unemployment	35.22 (1.06)	576	35.37 (1.00)	586	35.30 (0.73)	1,162	0.918
Inflation	103.60 (2.02)	576	100.06 (2.01)	586	101.82 (1.43)	1,162	0.216
Dollar blue	924.66 (4.82)	576	928.64 (4.61)	586	926.67 (3.33)	1,162	0.551
Official dollar	458.91 (8.68)	576	462.16 (9.03)	586	460.55 (6.27)	1,162	0.795

Notes: The table reports balance tests for respondents' reported current values of the economic indicators by accuracy-treatment status. Control refers to respondents who did not receive monetary incentives for accurate economic forecasts, and Treatment refers to respondents who did. Columns report means, standard errors in parentheses, and sample sizes. The final column reports p-values from two-sample t-tests of equality of means across treatment groups, allowing unequal variances. *, **, and *** indicate significance at the 10%, 5%, and 1% levels.

Table OA.5—: Massa Support and Reported Current Values

	Dollar blue	Official dollar	Inflation	Unemployment
Massa is favourite	-7.89 (8.91)	-6.48 (17.53)	-6.76* (3.95)	-4.59** (1.98)
Observations	1030	1030	1030	1030

Notes: The table reports coefficients from separate regressions of respondents' reported current values of each economic indicator on an indicator for whether Massa was the respondent's favorite candidate. The omitted group is respondents whose favorite candidate was Bullrich or Milei. Positive coefficients indicate that Massa supporters reported higher current values than other respondents, while negative coefficients indicate lower reported current values. Robust standard errors are in parentheses. *, **, and *** indicate significance at the 10%, 5%, and 1% levels.

Table OA.6—: Reported Current Values by Candidate Support

	Official dollar	Dollar blue	Inflation	Unemployment
Massa	458.13 (214.54)	923.60 (98.40)	99.18 (47.89)	30.77 (23.18)
Bullrich	466.90 (216.87)	923.73 (129.12)	102.34 (46.17)	35.45 (24.22)
Milei	445.34 (199.18)	929.63 (113.04)	103.97 (49.68)	35.38 (24.63)

Notes: The table reports mean reported current values of each economic indicator by respondents' favorite candidate. Standard deviations are in parentheses. Exchange-rate values are in ARS, and inflation and unemployment are reported in percentage points.

Table OA.7—: Other Beliefs and Mean Forecasts by Favorite Candidate

Panel A: Most important current issue in Argentina			
	Most important issue		
	Inflation	Unemployment	Both
Massa	55.36	0.60	44.05
Bullrich	45.42	0.68	53.90
Milei	53.62	2.12	44.27
Panel B: Most corrupt candidate			
	Most corrupt candidate		
	Massa	Bullrich	Milei
Massa	16.67	59.52	23.81
Bullrich	92.88	3.05	4.07
Milei	90.30	8.99	0.71
Panel C: Average official dollar			
	Official dollar		
	Massa	Bullrich	Milei
Massa	791.83	1021.19	1440.87
Bullrich	874.92	811.77	1081.20
Milei	899.70	879.09	895.45
Panel D: Average dollar blue			
	Dollar blue		
	Massa	Bullrich	Milei
Massa	1411.24	1635.48	2102.70
Bullrich	1651.08	1346.16	1645.22
Milei	1848.94	1600.48	1415.05
Panel E: Average inflation rate			
	Inflation		
	Massa	Bullrich	Milei
Massa	130.20	144.22	181.50
Bullrich	153.97	126.30	144.08
Milei	183.16	148.00	120.14
Panel F: Average unemployment rate			
	Unemployment		
	Massa	Bullrich	Milei
Massa	24.33	33.05	38.68
Bullrich	38.87	31.02	37.13
Milei	37.76	32.99	26.30

Notes: The table reports descriptive statistics by respondents' favorite candidate. Rows indicate the respondent's favorite candidate. Panels A and B report row percentages for the most important current issue and the candidate perceived as most corrupt. Panels C–F report mean candidate-contingent forecasts for each economic indicator, with columns indicating the candidate scenario under which the forecast was elicited. Exchange-rate forecasts are in ARS, and inflation and unemployment forecasts are in percentage points.

Table OA.8—: Forecast Dispersion by Favorite Candidate

Panel A: Standard deviation of official dollar			
	Official dollar		
	Massa	Bullrich	Milei
Massa	457.50	564.87	1026.93
Bullrich	517.05	469.75	812.20
Milei	586.05	557.56	704.96
Panel B: Standard deviation of dollar blue			
	Dollar blue		
	Massa	Bullrich	Milei
Massa	587.96	692.57	1440.19
Bullrich	714.87	505.99	998.93
Milei	816.61	689.00	819.28
Panel C: Standard deviation of inflation rate			
	Inflation		
	Massa	Bullrich	Milei
Massa	71.29	82.89	170.17
Bullrich	97.54	78.68	116.05
Milei	177.78	130.89	117.29
Panel D: Standard deviation of unemployment rate			
	Unemployment		
	Massa	Bullrich	Milei
Massa	21.67	25.04	27.31
Bullrich	29.31	23.21	26.95
Milei	29.78	26.48	22.92

Notes: The table reports standard deviations of candidate-contingent forecasts by respondents' favorite candidate. Rows indicate the respondent's favorite candidate, and columns indicate the candidate scenario under which the forecast was elicited. Exchange-rate forecasts are in ARS, and inflation and unemployment forecasts are in percentage points.

Table OA.9—: Kolmogorov–Smirnov Tests for Maximum Forecast Gaps

Comparison	Outcome	KS statistic	p-value	Control N	Treatment N
Accuracy treatment, no information	Official dollar	0.053	0.851	268	262
Accuracy treatment, no information	Dollar blue	0.152	0.005	268	261
Accuracy treatment, no information	Inflation	0.079	0.455	240	234
Accuracy treatment, no information	Unemployment	0.121	0.041	268	262
Information treatment, no accuracy	Official dollar	0.052	0.868	268	263
Information treatment, no accuracy	Dollar blue	0.106	0.103	268	263
Information treatment, no accuracy	Inflation	0.067	0.656	240	245
Information treatment, no accuracy	Unemployment	0.145	0.007	268	263
Both treatments vs. control	Official dollar	0.099	0.149	268	265
Both treatments vs. control	Dollar blue	0.090	0.227	268	263
Both treatments vs. control	Inflation	0.076	0.506	240	232
Both treatments vs. control	Unemployment	0.151	0.004	268	265

Notes: The table reports two-sample Kolmogorov–Smirnov tests comparing the distributions of maximum forecast gaps across treatment conditions. The maximum forecast gap is defined as the difference between the highest and lowest candidate-contingent forecast across the three candidates. The first panel compares the accuracy-treatment group with the pure control group among respondents who did not receive the information treatment. The second panel compares the information-treatment group with the pure control group among respondents who did not receive the accuracy treatment. The third panel compares respondents who received both treatments with the pure control group. The KS statistic is the maximum distance between the two empirical cumulative distribution functions.

Table OA.10—: Decomposition of Partisan Forecast Gaps

	Accuracy treatment, no information			Information treatment, no accuracy		
	(1) Favorite candidate	(2) Other candidates	(3) Gap	(4) Favorite candidate	(5) Other candidates	(6) Gap
Treatment coefficient	-0.039 (0.043)	-0.162*** (0.050)	-0.124*** (0.047)	-0.072* (0.042)	-0.163*** (0.050)	-0.091* (0.050)
Control mean	-0.237	0.118	0.355	-0.237	0.118	0.355
Observations	2011	2011	2011	2015	2015	2015

Notes: Gap is other candidates minus favorite candidate. All outcomes are z-scores.

Standard errors clustered by respondent. Stars denote p-values below 0.10, 0.05, and 0.01.

Notes: The table decomposes treatment effects on partisan forecast gaps into changes in forecasts under the respondent’s favorite candidate and changes in forecasts under the other two candidates. Forecasts are standardized by indicator using the mean and standard deviation in the pure control group, and higher values correspond to worse economic outcomes. “Other candidates” is the average forecast across the two non-favorite candidate scenarios, and “Gap” is defined as other candidates minus favorite candidate. Columns (1)–(3) compare the accuracy-treatment group with the control group among respondents who did not receive information. Columns (4)–(6) compare the information-treatment group with the control group among respondents who did not receive accuracy incentives. All regressions include indicator fixed effects. Standard errors clustered at the respondent level are in parentheses. *, **, and *** indicate significance at the 10%, 5%, and 1% levels.

Table OA.11—: Treatment Effects on the Other-Candidate Forecast Penalty

	(1) Pooled z-score	(2) Official dollar	(3) Dollar blue	(4) Inflation	(5) Unemployment
Other candidate	0.377*** (0.040)	245.335*** (48.958)	442.158*** (74.628)	39.869*** (9.625)	11.191*** (1.476)
Accuracy $\times$ other candidate	-0.117** (0.050)	-91.608 (64.675)	-112.674 (88.791)	-10.484 (11.724)	-3.846** (1.840)
Information $\times$ other candidate	-0.099* (0.053)	-38.184 (63.750)	-111.590 (102.880)	-5.736 (12.603)	-4.752** (2.042)
Accuracy $\times$ information $\times$ other candidate	0.103 (0.067)	39.349 (87.203)	53.717 (122.608)	3.801 (15.434)	7.054*** (2.668)
Respondents	1030	1030	1027	928	1030
Respondent FE	Yes	Yes	Yes	Yes	Yes
Candidate FE	Yes	Yes	Yes	Yes	Yes
Indicator FE	Yes	No	No	No	No
Observations	12043	3090	3079	2784	3090

Notes: The table reports regressions of candidate-contingent forecasts on an indicator for whether the forecast refers to a candidate other than the respondent's favorite candidate, interactions between this indicator and the two treatments, respondent fixed effects, and candidate fixed effects. Column (1) pools all four indicators after standardizing forecasts using the mean and standard deviation in the pure control group, and includes indicator fixed effects. Columns (2)–(5) report indicator-specific regressions in original units. Higher values correspond to worse economic outcomes for all indicators. The coefficient on “Other candidate” captures the forecast penalty assigned to non-favorite candidates in the pure control group. Negative interaction coefficients indicate that the treatment reduces this other-candidate penalty. Standard errors clustered at the respondent level are in parentheses. *, **, and *** indicate significance at the 10%, 5%, and 1% levels.

Table OA.12—: Treatment Effects on Forecast Distances

Panel A. Milei forecast vs. realized future value				
	(1) Official dollar	(2) Dollar blue	(3) Inflation	(4) Unemployment
Accuracy treatment, no information	18.28 (53.63)	-109.85 (77.44)	-10.16 (7.55)	-3.89* (2.08)
Avg. Control	545.13	761.63	182.29	30.83
Observations	517	515	463	517
Information treatment, no accuracy	55.51 (57.40)	61.57 (97.92)	-11.52 (8.52)	-9.00*** (2.19)
Avg. Control	545.13	761.63	182.29	30.83
Observations	515	515	470	515
Panel B. Average non-Milei contingent forecasts vs. realized future value				
	(1) Official dollar	(2) Dollar blue	(3) Inflation	(4) Unemployment
Accuracy treatment, no information	-93.14*** (34.49)	-83.93 (55.31)	-7.56 (8.10)	-5.55** (2.17)
Avg. Control	448.21	731.86	164.58	34.11
Observations	517	516	463	517
Information treatment, no accuracy	-118.22*** (32.04)	-60.79 (55.40)	-16.37** (7.98)	-11.73*** (2.26)
Avg. Control	448.21	731.86	164.58	34.11
Observations	515	515	470	515
Panel C. Milei forecast vs. current value shown in information treatment				
	(1) Official dollar	(2) Dollar blue	(3) Inflation	(4) Unemployment
Accuracy treatment, no information	4.96 (67.26)	-114.77 (78.02)	-17.05** (8.20)	-3.86* (2.11)
Avg. Control	713.48	808.55	80.99	32.02
Observations	517	515	463	517
Information treatment, no accuracy	62.81 (70.78)	59.68 (98.50)	1.34 (10.23)	-9.24*** (2.22)
Avg. Control	713.48	808.55	80.99	32.02
Observations	515	515	470	515
Panel D. Average contingent forecasts vs. current value shown in information treatment				
	(1) Official dollar	(2) Dollar blue	(3) Inflation	(4) Unemployment
Accuracy treatment, no information	-58.69 (48.77)	-103.63** (52.38)	-17.98** (7.79)	-4.95** (2.04)
Avg. Control	643.27	774.07	77.73	34.23
Observations	517	516	463	517
Information treatment, no accuracy	-50.24 (47.88)	-24.34 (55.47)	-9.69 (7.77)	-11.08*** (2.13)
Avg. Control	643.27	774.07	77.73	34.23
Observations	515	515	470	515

Notes: The table reports treatment effects on absolute forecast distances. Negative coefficients indicate that forecasts are closer to the benchmark value. Panel A uses the distance between the respondent's Milei-contingent forecast and the realized future value, since Milei won the election. Panel B uses the distance between the average of the Massa- and Bullrich-contingent forecasts and the realized future value. Panel C uses the distance between the Milei-contingent forecast and the current value shown in the information treatment. Panel D uses the distance between the average of all three candidate-contingent forecasts and the current value shown in the information treatment. The realized future values are 876 ARS for the official dollar, 1005 ARS for the blue dollar, 287.9% for annual inflation, and 7.7% for unemployment. The current values shown in the information treatment are 357 ARS, 940 ARS, 124.4%, and 6.2%. All regressions control for favorite-candidate fixed effects. Robust standard errors are in parentheses. *, **, and *** indicate significance at the 10%, 5%, and 1% levels.

Table OA.13—: Beliefs about Blue-Dollar Overestimation

	(1) Only Massa	(2) Only Bullrich	(3) Only Milei
Preferred candidate (no accuracy, no information)	0.01 (0.09)	-0.23*** (0.08)	0.46*** (0.05)
Observations	150.00	213.00	420.00
Preferred candidate (accuracy only)	-0.17*** (0.06)	-0.14* (0.08)	0.32*** (0.06)
Observations	123.00	189.00	456.00
Preferred candidate (information only)	-0.10 (0.10)	-0.14** (0.06)	0.31*** (0.06)
Observations	108.00	249.00	405.00
Preferred candidate (accuracy and information)	-0.15* (0.07)	-0.19*** (0.07)	0.37*** (0.06)
Observations	123.00	234.00	420.00
Individual FE	Yes	Yes	Yes

Notes: The table reports within-respondent regressions for beliefs about the direction of blue-dollar forecast errors. The outcome equals one if the respondent expects to have overestimated the blue dollar, that is, if they report that the realized blue dollar would be lower than their forecast conditional on their forecast being wrong. Columns restrict the sample to respondents whose favorite candidate is Massa, Bullrich, or Milei. The coefficient on “Preferred candidate” compares the respondent’s preferred-candidate scenario with the two non-preferred candidate scenarios within the same respondent. Each row reports this comparison for a different treatment condition. All regressions include respondent fixed effects. Standard errors clustered at the respondent level are in parentheses. *, **, and *** indicate significance at the 10%, 5%, and 1% levels.

Table OA.14—: Accuracy of Expected Error Direction

	(1) Pool data	(2) Only Massa	(3) Only Bullrich	(4) Only Milei
Correct direction (no accuracy, no information)	0.34 (0.03)	0.18 (0.05)	0.32 (0.06)	0.42 (0.04)
p-value: share = 0.50	0.000	0.000	0.002	0.063
Observations	268	50	71	140
Correct direction (accuracy only)	0.37 (0.03)	0.32 (0.07)	0.32 (0.06)	0.41 (0.04)
p-value: share = 0.50	0.000	0.017	0.003	0.033
Observations	260	41	63	150
Correct direction (information only)	0.39 (0.03)	0.14 (0.06)	0.36 (0.05)	0.47 (0.04)
p-value: share = 0.50	0.000	0.000	0.011	0.549
Observations	263	36	83	135
Correct direction (accuracy and information)	0.38 (0.03)	0.32 (0.07)	0.34 (0.05)	0.41 (0.04)
p-value: share = 0.50	0.000	0.017	0.005	0.042
Observations	263	41	76	140

Notes: The table reports the share of respondents whose expected direction of error for the Milei-contingent blue-dollar forecast matches the realized direction of error. Because Milei won the election, only the Milei-contingent forecast can be compared with the realized future value. A response is coded as correct if the respondent expected to overestimate and their forecast was above the realized value, or if they expected to underestimate and their forecast was below the realized value. Forecasts exactly equal to the realized value are excluded because the question was conditional on the forecast being wrong. Columns report the pooled sample and subsamples by favorite candidate. Each row corresponds to a treatment condition. Standard errors are in parentheses. The p-value tests whether the share equals 0.50.

Table OA.15—: Expected Error Direction and the Most-Different Forecast

	(1) Only Massa	(2) Only Bullrich	(3) Only Milei
Toward most different forecast (no accuracy, no information)	0.74 (0.07)	0.84 (0.05)	0.41 (0.05)
p-value: share = 0.50	0.001	0.000	0.044
Observations	42	56	120
Toward most different forecast (accuracy only)	0.76 (0.08)	0.71 (0.06)	0.37 (0.04)
p-value: share = 0.50	0.002	0.001	0.002
Observations	33	56	131
Toward most different forecast (information only)	0.69 (0.08)	0.77 (0.05)	0.45 (0.05)
p-value: share = 0.50	0.032	0.000	0.271
Observations	32	61	118
Toward most different forecast (accuracy and information)	0.64 (0.08)	0.79 (0.05)	0.40 (0.04)
p-value: share = 0.50	0.078	0.000	0.017
Observations	39	63	129

Notes: The table reports the share of respondents whose expected error direction for their preferred-candidate blue-dollar forecast points toward the most different of their two non-preferred candidate-contingent blue-dollar forecasts. For each respondent, the most different forecast is the non-preferred candidate forecast farthest from the preferred-candidate forecast. The expected direction points toward that forecast if the respondent expects the realized blue dollar to be higher when the most different forecast is higher, or lower when the most different forecast is lower. Cases with ambiguous ties in opposite directions are excluded. Columns restrict the sample by favorite candidate, and rows correspond to treatment conditions. Standard errors are in parentheses. The p-value tests whether the share equals 0.50.

Table OA.16—: Accuracy Incentives and Forecast Gaps, Winsorized Forecast

	(1) Official dollar	(2) Dollar blue	(3) Inflation	(4) Unemployment
Maximum difference	-1.51 (39.01)	-102.86** (47.72)	-2.47 (5.30)	-4.40*** (1.46)
Avg. Control	450.54 (27.17)	669.77 (35.36)	50.38 (3.75)	19.20 (1.15)
Observations	530	529	474	530
Massa and Milei only	-89.40 (61.87)	-66.54 (70.94)	1.14 (7.35)	-4.07* (2.09)
Avg. Control	175.91 (42.65)	453.24 (52.16)	38.91 (5.24)	14.24 (1.67)
Observations	383	380	341	383
Massa and Bullrich only	-21.68 (47.39)	-142.68** (58.55)	-7.86 (5.06)	-1.79 (1.85)
Avg. Control	120.02 (38.79)	335.52 (45.33)	22.85 (3.90)	9.32 (1.50)
Observations	225	225	204	225
Milei and Bullrich only	-109.34*** (42.24)	-95.31* (52.23)	-2.31 (4.85)	-3.33** (1.42)
Avg. Control	122.65 (29.57)	291.34 (41.13)	20.23 (3.68)	7.92 (1.10)
Observations	426	424	381	426
Least - most favourite	-68.08 (46.71)	-100.36* (55.22)	-2.31 (5.64)	-4.17** (1.65)
Avg. Control	154.51 (32.14)	458.19 (40.51)	35.46 (3.87)	13.75 (1.31)
Observations	517	515	463	517

Notes: The table reports coefficients from regressions of forecast-gap measures on the accuracy-treatment indicator, restricted to respondents who did not receive the information treatment. Candidate-level forecasts are winsorized at the 5th and 95th percentiles before forecast gaps are computed. The main coefficients are in the first row, where the outcome is the maximum difference across the three candidate-contingent forecasts for each indicator. Robust standard errors are in parentheses. The control-group mean is reported for each outcome. *, **, and *** indicate significance at the 10%, 5%, and 1% levels.

Table OA.17—: Information Provision and Forecast Gaps, Winsorized Forecasts

	(1) Official dollar	(2) Dollar blue	(3) Inflation	(4) Unemployment
Maximum difference	42.79 (42.30)	-79.31 (49.09)	2.36 (5.33)	-3.69** (1.62)
Avg. Control	450.54 (27.17)	669.77 (35.36)	50.38 (3.75)	19.20 (1.15)
Observations	531	531	485	531
Massa and Milei only	-48.19 (67.23)	-58.60 (73.42)	4.15 (7.53)	-4.66* (2.39)
Avg. Control	175.91 (42.65)	453.24 (52.16)	38.91 (5.24)	14.24 (1.67)
Observations	361	361	326	361
Massa and Bullrich only	-10.74 (55.36)	-126.34* (64.97)	2.82 (5.73)	-2.92 (2.00)
Avg. Control	120.02 (38.79)	335.52 (45.33)	22.85 (3.90)	9.32 (1.50)
Observations	240	240	224	240
Milei and Bullrich only	-39.28 (45.81)	-124.87** (55.26)	-4.22 (5.21)	-3.42** (1.55)
Avg. Control	122.65 (29.57)	291.34 (41.13)	20.23 (3.68)	7.92 (1.10)
Observations	429	429	390	429
Least - most favourite	-54.04 (50.33)	-148.51*** (57.04)	0.19 (5.55)	-5.18*** (1.81)
Avg. Control	154.51 (32.14)	458.19 (40.51)	35.46 (3.87)	13.75 (1.31)
Observations	515	515	470	515

Notes: The table reports coefficients from regressions of forecast-gap measures on the information-treatment indicator, restricted to respondents who did not receive the accuracy treatment. Candidate-level forecasts are winsorized at the 5th and 95th percentiles before forecast gaps are computed. The main coefficients are in the first row, where the outcome is the maximum difference across the three candidate-contingent forecasts for each indicator. Robust standard errors are in parentheses. The control-group mean is reported for each outcome. *, **, and *** indicate significance at the 10%, 5%, and 1% levels.

Table OA.18—: Attention-Check Failure Rates

Attention check	Group	Failed	Non-missing responses	Failure rate
Failed at least one	Overall	107	856	12.5%
Failed at least one	No accuracy, no information	33	218	15.1%
Failed at least one	Accuracy only	22	212	10.4%
Failed at least one	Information only	26	214	12.1%
Failed at least one	Accuracy + information	26	212	12.3%
Failed both	Overall	47	838	5.6%
Failed both	No accuracy, no information	9	214	4.2%
Failed both	Accuracy only	10	208	4.8%
Failed both	Information only	12	212	5.7%
Failed both	Accuracy + information	16	204	7.8%

Notes: The table reports attention-check failure rates overall and by 2x2 treatment condition. Respondents pass each attention check if they selected the instructed response, “Muy de acuerdo.” “Failed at least one” equals one if the respondent failed either attention check, and “Failed both” equals one if the respondent failed both checks. Non-missing responses report the number of respondents with observed answers for the relevant attention-check measure. Failure rates are calculated as failed responses divided by non-missing responses.

Table OA.19—: Accuracy Incentives and Forecast Gaps, Attention-Check Sample

	(1) Official dollar	(2) Dollar blue	(3) Inflation	(4) Unemployment
Maximum difference	3.37 (70.14)	-100.21 (91.04)	-14.48 (13.72)	-3.75** (1.70)
Avg. Control	523.74 (50.15)	785.26 (73.73)	74.45 (10.99)	18.08 (1.30)
Observations	369	368	328	369
Massa and Milei only	3.37 (99.97)	-83.75 (124.86)	4.82 (17.69)	-2.30 (2.40)
Avg. Control	146.07 (69.58)	598.47 (102.32)	50.63 (14.01)	12.55 (1.84)
Observations	279	276	244	279
Massa and Bullrich only	10.79 (64.87)	-3.65 (82.45)	-10.08* (5.18)	-2.51 (1.96)
Avg. Control	72.92 (54.76)	306.22 (46.63)	22.56 (3.87)	8.23 (1.58)
Observations	160	160	147	160
Milei and Bullrich only	-101.37 (67.68)	16.88 (70.36)	-6.89 (12.86)	-2.46 (1.76)
Avg. Control	114.04 (50.50)	259.80 (52.31)	29.96 (11.87)	7.98 (1.33)
Observations	299	297	265	299
Least - most favourite	-13.63 (78.28)	-42.62 (93.58)	1.70 (12.31)	-3.11* (1.89)
Avg. Control	146.64 (54.99)	543.05 (72.84)	42.81 (9.01)	12.48 (1.42)
Observations	369	367	328	369

Notes: The table reports coefficients from regressions of forecast-gap measures on the accuracy-treatment indicator, restricted to respondents who did not receive the information treatment and who passed both attention checks. The main coefficients are in the first row, where the outcome is the maximum difference across the three candidate-contingent forecasts for each indicator. Robust standard errors are in parentheses. The control-group mean is reported for each outcome. *, **, and *** indicate significance at the 10%, 5%, and 1% levels.

Table OA.20—: Information Provision and Forecast Gaps, Attention-Check Sample

	(1) Official dollar	(2) Dollar blue	(3) Inflation	(4) Unemployment
Maximum difference	92.81 (79.06)	41.23 (115.08)	-5.52 (13.95)	-3.18* (1.88)
Avg. Control	523.74 (50.15)	785.26 (73.73)	74.45 (10.99)	18.08 (1.30)
Observations	368	368	336	368
Massa and Milei only	15.98 (109.43)	9.65 (155.76)	6.36 (18.68)	-4.68* (2.77)
Avg. Control	146.07 (69.58)	598.47 (102.32)	50.63 (14.01)	12.55 (1.84)
Observations	257	257	232	257
Massa and Bullrich only	13.08 (74.99)	-44.03 (80.29)	1.93 (6.15)	-1.85 (2.20)
Avg. Control	72.92 (54.76)	306.22 (46.63)	22.56 (3.87)	8.23 (1.58)
Observations	181	181	169	181
Milei and Bullrich only	-29.59 (75.82)	-56.96 (85.61)	-14.11 (12.96)	-5.51*** (1.85)
Avg. Control	114.04 (50.50)	259.80 (52.31)	29.96 (11.87)	7.98 (1.33)
Observations	298	298	271	298
Least - most favourite	-17.21 (83.17)	-82.89 (111.49)	-0.76 (11.52)	-5.30** (2.06)
Avg. Control	146.64 (54.99)	543.05 (72.84)	42.81 (9.01)	12.48 (1.42)
Observations	368	368	336	368

Notes: The table reports coefficients from regressions of forecast-gap measures on the information-treatment indicator, restricted to respondents who did not receive the accuracy treatment and who passed both attention checks. The main coefficients are in the first row, where the outcome is the maximum difference across the three candidate-contingent forecasts for each indicator. Robust standard errors are in parentheses. The control-group mean is reported for each outcome. *, **, and *** indicate significance at the 10%, 5%, and 1% levels.

Table OA.21—: Incentives and Placebo Forecasts

	MERVAL	Bitcoin	Temperature	Football
Incentive treatment	-0.03 (0.03)	-0.04 (0.03)	-1.61*** (0.43)	-0.04 (0.03)
Control mean	1.35	1.33	32.13	0.53
Observations	1043	1043	1035	979

Notes: The table reports coefficients from separate regressions of each placebo forecast outcome on an indicator for receiving monetary incentives for the placebo questions. MERVAL and Bitcoin are coded using the original directional forecast responses. Temperature is reported in degrees Celsius. Football is an indicator equal to one when the respondent's forecast for the best-performing team is congruent with their stated favorite team. Robust standard errors are in parentheses. The control mean is reported for respondents who did not receive incentives for the placebo questions. *, **, and *** indicate significance at the 10%, 5%, and 1% levels.

Motivated Forecasts - Elections in Argentina 2023 (#147198)

Author(s)

This pre-registration is currently anonymous to enable blind peer-review.
It has one author.

Pre-registered on: 10/15/2023 11:43 AM (PT)

1) Have any data been collected for this study already?

No, no data have been collected for this study yet.

2) What's the main question being asked or hypothesis being tested in this study?

- H1: Do accuracy incentives reduce the political gap in forecasts?
- H2a: Informing subjects about the current levels of economic indicators reduces variance in the forecasts?
- H2b: Does the information treatment reduce the political gap?
- H3a: Do they correctly guess in which direction they might be missforecasting?
- H3b: Do they guess they will missforecast in the direction of their forecast for the most different candidate?
- H4a: Do the incentives have a stronger effect on the forecasts for own candidate or on the forecasts regarding other candidates?
- H4b: Do the incentives have a stronger effect on the forecasts for the candidate they expect to win?
- H5a: Are the effects different by literacy?
- H5b: Are the effects different by their prior knowledge of the current levels of economic indicators?
- H6a: Do people who spend more time in each question get more accurate responses to the current levels?
- H6b: Are the treatment effects stronger for people who spend more time in the forecasting questions?
- H7: Do people agree on which candidate would be better for each specific issue regardless of their preferences for the winning candidate?
- H8: Do accuracy incentives improve people's ability to forecast economic indicators?

3) Describe the key dependent variable(s) specifying how they will be measured.

Political gap: We will create a "Forecast Gap" score, which refers to the maximal difference between the forecasts of the three candidates, for each of the four questions (inflation, unemployment, dollar blue, official dollar). We expect to find a smaller gap in the incentivised condition. Similarly, we expect to find a smaller gap with the information treatment.

Gap in other topics: We will create a binary variable that takes the value 1 if the forecasts is congruent with their football team and 0 if not (e.g., if a Racing supporter forecasts that Racing will do better, we will code it as 1, and 0 if not). The teams have a classic rival (Boca/River and Racing/Independiente). We will also code the forecasts of the worst team (e.g., if a Racing supporter forecasts that Independiente will be the worst, we will code it as 1, and 0 if not). We will use the responses to the temperature, Merval and Bitcoin price as dependent variables. We expect to see motivated forecasts in the football questions but not in the other three (temperature, Merval and Bitcoin).

Accuracy will be measured by the difference between the response or the forecast and the real value.

4) How many and which conditions will participants be assigned to?

2x2 factorial design:

- 1) Condition A: Incentives for accuracy in inflation, unemployment and price of the informal (blue) and official dollar. In the incentivised questions: For every correct answer, participants will increase their chances of winning a prize in a lottery. No incentives for other questions.
- 2) Condition B: No incentives for accuracy in inflation, unemployment and price of the informal (blue) and official dollar. Incentives for accuracy in questions regarding football, temperature, Merval and Bitcoins. In the incentivised questions: For every correct answer, participants will increase their chances of winning a prize in a lottery.

In each condition participants will be further divided into two groups: one of them will be informed about the correct levels of inflation, unemployment and the price of the dollar (blue and official), and the other will not receive any information.

5) Specify exactly which analyses you will conduct to examine the main question/hypothesis.

The main analysis will be conducted using regression analysis.

Primary Analyses:

- 1. We will test if accuracy incentives reduce the political gap in forecasts.
- 2. We will test if informing subjects about the current levels of economic indicators reduces variance in the forecasts.
- 3. We will test if the information treatment reduce the political gap in forecasts.
- 4. We will test whether the incentives have a stronger effect on the forecasts for own candidate or on the forecasts regarding other candidates.

5. We will test if the incentives have a stronger effect on the forecasts for the candidate they expect to win.
6. We will test whether the main effects are different by literacy and prior knowledge about the economic indicators.
7. We will test these hypotheses separately for each group of supporters (Bullrich, Massa and Milei).

Secondary Analyses:

1. We will test if there is positive effect of the incentives on the forecast accuracy.
2. We will test if the participants prefer the candidates who they expect will perform better in their preferred issue.
3. We will test if conditional on the preferred candidate, people who think inflation/unemployment is more important, forecast better outcomes than their counterparts.
4. We will test if people who spend more time in each question get more accurate responses to the current levels.
5. We will test if the treatment effects are stronger for people who spend more time in the forecasting questions.
6. We will test if people are able to predict the direction of their forecasting error.
7. We will test if the direction of the forecasting error is in the direction of the most different candidate.

Exploratory Analyses:

1. We will explore if the difference in forecasts made about different candidates are larger when "made about the same candidate by different group of supporters" or when "made about different candidates by the same group of supporters". This is to understand whether people have different views about the future performance of the candidates vs their preferences for certain issues.
2. We will explore which supporters are more informed about the current levels.
3. We will explore if people agree on which candidate would be better for each specific issue regardless of their preferences for the winning candidate.

6) Describe exactly how outliers will be defined and handled, and your precise rule(s) for excluding observations.

As a robustness test we will exclude observations that did not answer the two attention checks correctly.

7) How many observations will be collected or what will determine sample size? No need to justify decision, but be precise about exactly how the number will be determined.

We will send invitations to 2400 students and alumni of the Universidad Nacional del Nordeste (Argentina) who registered to participate in online experiments. We are time constrained by the fact that the elections will take place on October 22nd. We will send the last reminder on October 20th and will keep the survey open until the end of October 21st. Following Rathje et al (2023), we aim for at least 500 participants.

Rathje, S., Roozenbeek, J., Van Bavel, J. J., & van der Linden, S. (2023). Accuracy and social motivations shape judgements of (mis) information. *Nature human behaviour*, 1-12.

8) Anything else you would like to pre-register? (e.g., secondary analyses, variables collected for exploratory purposes, unusual analyses planned?)

Nothing else to pre-register.

BUNDLE

This pre-registration is part of a set of similar and/or related pre-registrations sharing at least one author. When one of these pre-registrations was shared by an author, the rest were shared automatically. Links to other pre-registration(s), appear below:

#147200 - <https://aspredicted.org/h7tj-fh3j.pdf> - Title: 'Conscription and its effects on the next generation'

IMPORTANTE: La información correcta sobre las preguntas anteriores es la siguiente:

Nivel de desocupación en el 2do trimestre de 2023: **6,2%**

Nivel de inflación interanual en Agosto de 2023: **124,4%**

Valor del Dólar Blue al 9 de Octubre de 2023: **\$940**

Valor del Dólar Oficial al 9 de Octubre de 2023: **\$357**

Instructions with English Translations

Screenshot 1

A continuación nos gustaría saber que pensás sobre los siguientes indicadores

Cuál es el nivel de desocupación actual? En porcentaje

0 10 20 30 40 50 60 70 80 90 100

Cual es el nivel de desempleo actual?



Qué tan seguro/a estás sobre tu respuesta? (0 es nada seguro/a y 100 es completamente seguro/a)

0 10 20 30 40 50 60 70 80 90 100

Confianza en mi respuesta



English translation

Next, we would like to know what you think about the following indicators.

What is the current unemployment rate? As a percentage.

What is the current unemployment level?

How sure are you about your answer? (0 is not at all sure and 100 is completely sure.)

Confidence in my answer.

Screenshot 2

Cuál es el nivel de inflación interanual en Agosto de 2023? En porcentaje ANUAL

Qué tan seguro/a estás sobre tu respuesta? (0 es nada seguro/a y 100 es completamente seguro/a)

0 10 20 30 40 50 60 70 80 90 100

Confianza en mi respuesta



English translation

What is the year-on-year inflation rate in August 2023? As an ANNUAL percentage.

How sure are you about your answer? (0 is not at all sure and 100 is completely sure.)

Confidence in my answer.

Screenshot 3

Cuál fue el valor del dólar blue el Lunes 9 de Octubre de 2023?

Qué tan seguro/a estás sobre tu respuesta? (0 es nada seguro/a y 100 es completamente seguro/a)

0 10 20 30 40 50 60 70 80 90 100

Confianza en mi respuesta



English translation

What was the value of the blue dollar on Monday, October 9, 2023?

How sure are you about your answer? (0 is not at all sure and 100 is completely sure.)

Confidence in my answer.

Screenshot 4

Cuál fue el valor del dólar oficial el Lunes 9 de Octubre de 2023?

Qué tan seguro/a estás sobre tu respuesta? (0 es nada seguro/a y 100 es completamente seguro/a)

0 10 20 30 40 50 60 70 80 90 100

Confianza en mi respuesta



English translation

What was the value of the official dollar on Monday, October 9, 2023?

How sure are you about your answer? (0 is not at all sure and 100 is completely sure.)

Confidence in my answer.

Screenshot 5 (Information treatment)

IMPORTANTE: La información correcta sobre las preguntas anteriores es la siguiente:

Nivel de desocupación en el 2do trimestre de 2023: **6,2%**

Nivel de inflación interanual en Agosto de 2023: **124,4%**

Valor del Dólar Blue al 9 de Octubre de 2023: **\$940**

Valor del Dólar Oficial al 9 de Octubre de 2023: **\$357**

English translation

IMPORTANT: The correct information about the previous questions is as follows:

Unemployment rate in the 2nd quarter of 2023: 6.2%

Year-on-year inflation rate in August 2023: 124.4%

Value of the blue dollar on October 9, 2023: \$940

Value of the official dollar on October 9, 2023: \$357

Screenshot 6 (Accuracy treatment)

IMPORTANTE!

Ahora te pedimos que hagas un pronóstico de las siguientes preguntas. En esta sección podrás aumentar la probabilidad de ganar uno de los **30 premios de 100 dólares** en efectivo. Con cada respuesta correcta obtendrás **10 boletos** más para la lotería. Se considerarán correctas a las respuestas que no se alejen más de un 5% por encima o por debajo del valor real. Formalmente, si tu pronóstico es x , tu respuesta será correcta si el valor real está entre $x+0.05x$ y $x-0.05x$.

English translation

IMPORTANT!

Now we ask you to make forecasts for the following questions. In this section you will be able to increase the probability of winning one of the 30 cash prizes of 100 US dollars.

For each correct answer you will receive 10 more lottery tickets.

Answers that are no more than 5% above or below the real value will be considered correct.

Formally, if your forecast is x , your answer will be correct if the real value is between $x - 0.05x$ and $x + 0.05x$.

Screenshot 7

Cuál será el nivel de inflación interanual en Marzo de 2024 en cada caso? Porcentaje

Si gana Massa	
Si gana Bullrich	
Si gana Milei	

English translation

What will the year-on-year inflation rate be in March 2024 in each case? Percentage.

If Massa wins.

If Bullrich wins.

If Milei wins.

Screenshot 8

Cuál será el nivel de desocupación en el primer trimestre de 2024 en cada caso?

Porcentaje

0 10 20 30 40 50 60 70 80 90 100

Si gana Massa



Si gana Bullrich



Si gana Milei



English translation

What will the unemployment rate be in the first quarter of 2024 in each case? Percentage.

If Massa wins.

If Bullrich wins.

If Milei wins.

Screenshot 9

Cuál será el precio del dólar blue al primero de abril de 2024 en cada caso?

Si gana Massa

Si gana Bullrich

Si gana Milei

English translation

What will the price of the blue dollar be on April 1, 2024 in each case?

If Massa wins.

If Bullrich wins.

If Milei wins.

Screenshot 10

Cuál será el precio del dólar oficial al primero de abril de 2024 en cada caso?

Si gana Milei

Si gana Bullrich

Si gana Massa

English translation

What will the price of the official dollar be on April 1, 2024 in each case?

If Milei wins.

If Bullrich wins.

If Massa wins.

Screenshot 11

En caso de que tu predicción sobre el precio del dólar no sea correcta, que es más probable?

	Que el precio del dólar sea más bajo que tu predicción	Que el precio del dólar sea más alto que tu predicción
Bullrich	☐	☐
Milei	☐	☐
Massa	☐	☐

English translation

If your prediction about the dollar price is not correct, which is more likely?
Column options: The dollar price will be lower than your prediction; the dollar price will be higher than your prediction.
Rows: Bullrich, Milei, Massa.