

Public Approval Rating Updates for President Prabowo Using a Bayesian Dirichlet-Multinomial Hybrid Model

Anwar^{1,a,*}; Andries Riesfandhy^{2,b}

¹Faculty of Engineering, Informatics Program, Universitas Muhammadiyah, Palu, Indonesia.

²Faculty of Politics and Law, Political Science Program, Universitas Andi Sudirman, Bone, Indonesia.

^a Anwartalifoundation@gmail.com; ^b Politea.etos06@gmail.com

* Corresponding author

Abstract

Conventional estimation of presidential approval ratings relies on face-to-face surveys that excel in population representativeness (probability sampling) but suffer from significant time lag between data collection and publication. Conversely, social media monitoring (SMM) provides real-time data but is often biased toward active digital platform demographics. This study proposes a hybrid model based on Bayesian Dirichlet-Multinomial (BDM) to integrate survey data as prior belief with SMM data as likelihood to generate more dynamic and accurate posterior estimates. Prior data were drawn from the Indikator Politik Indonesia national survey of October 2025 (n=1,220, MoE $\pm 2.9\%$), indicating a public approval rate of 77.7%. Update data were collected from the X (Twitter) platform during November 2025 through crawling and scraping techniques, yielding 2.5 million raw data points processed through Named Entity Recognition (NER), text normalization, and Transformer-based classification (IndoBERT v2.0). After deduplication, 310,000 unique accounts were retained with 120 million total impressions. Sentiment distribution showed 61% positive, 15% negative, and 24% neutral. Bayesian updating produced a corrected public approval estimate of 68.4% (a 9.3 percentage point decline from the survey prior), detecting pockets of dissatisfaction on food price issues underrepresented in conventional surveys. Model validation using Leave-One-Out Cross-Validation (LOO-CV) produced an Expected Log-Pointwise Predictive Density (ELPD) of -127.4, superior to the single-survey baseline model (-148.9). This model offers a new framework as a public opinion early warning system that is responsive to current issue dynamics without sacrificing the statistical validity of traditional survey methods.

Keywords : Bayesian Dirichlet-Multinomial; Public Opinion Survey; Social Media Monitoring; Sentiment Analysis; IndoBERT; Indonesia.

1. Introduction

Public satisfaction with the performance of a head of state is the primary barometer of governmental legitimacy in a representative democratic system. In Indonesia, systematic measurement of presidential approval is conducted by survey institutions such as Indikator Politik Indonesia, Lembaga Survei Indonesia (LSI), and Saiful Mujani Research and Consulting (SMRC). These surveys employ face-to-face interview methods with probabilistic sampling

techniques that are theoretically capable of representing the national voting-age population (Dahl, 2015; Norris, 2011).

However, conventional survey methods face increasingly evident structural challenges in the digital era. First, there is a significant time lag between the data collection period and the moment of publication; a survey conducted in September may only be released in October or even December, while public opinion can shift drastically within days in response to controversial policies, viral events, or sudden crises (Pasek, 2018; Sanders et al., 2019). Second, the operational costs of national-scale face-to-face surveys are very high, limiting measurement frequency. Third, surveys are susceptible to social desirability bias — the tendency of respondents to give answers considered socially "safe" rather than their true opinions, particularly on sensitive political topics (Holbrook & Krosnick, 2010; Tourangeau & Yan, 2007).

On the other hand, social media platforms — especially X (formerly Twitter) — have evolved into a massive and unfiltered space for public opinion expression. Data from We Are Social (2025) shows that active social media users in Indonesia have reached 167 million people (59.4% of the total population), with X serving as the primary political discussion platform among opinion leaders, journalists, academics, and activists. This platform provides real-time sentiment data at volumes far exceeding the capacity of conventional surveys. Nevertheless, social media data carries serious methodological weaknesses: demographic bias skewed toward younger, highly educated urban groups (urban digital divide), the proliferation of bot accounts that can amplify fake sentiment, and echo chamber phenomena that distort opinion representation (Mellon & Prosser, 2017; Barbera, 2015; Colleoni et al., 2014).

Previous studies have attempted to use social media data as predictors of public opinion polls, with mixed results. (O'Connor et al. 2010) found correlations between Twitter sentiment and traditional polling in the United States, though with unstable fluctuations. (Ceron et al. 2014) in the Italian context found that social media sentiment aggregation can predict election results with adequate accuracy. However, in the Indonesian context — which has distinct digital ecosystem characteristics including high use of informal language, regional language mixing, and unique local issue dynamics — these models cannot be directly applied without methodological adaptation (Kurniawan & Louvan, 2022; Koto et al., 2020).

This study proposes a novel approach by integrating both data sources through a Bayesian statistical framework, specifically the Dirichlet-Multinomial model. In the Bayesian paradigm, valid prior knowledge (from surveys) is encoded as a prior distribution, then systematically updated using new evidence from social media as the likelihood, producing a posterior distribution that optimally combines the strengths of both sources (Gelman et al., 2013). This framework explicitly accommodates uncertainty from both data sources and enables theoretically and empirically grounded relative weighting.

Specifically, this study aims to: (1) develop technical procedures for integrating face-to-face survey data from Indikator Politik Indonesia (October 2025) with social media monitoring data (November 2025) using the BDM model; (2) analyze the dynamics of change in public approval estimates for the performance of President Prabowo Subianto following Bayesian updating; (3) identify the main issues driving sentiment shifts during the observation period; and (4) validate the hybrid model's performance against baseline approaches using standardized statistical metrics. This study thereby contributes to the development of a more responsive and cost-efficient public opinion measurement methodology, relevant to political science researchers, survey institutions, and policymakers who require near-real-time public opinion information.

2. Method

This study employs a computational quantitative approach that combines public opinion survey methods with big data social media analysis within a Bayesian inference framework. The research design is integrative-sequential, where data from two sources are collected at different yet temporally related periods.

2.1 Prior Data: Indikator Politik Indonesia Survey (October 2025)

Prior data were drawn from a national survey conducted by Indikator Politik Indonesia during the period of October 20–27, 2025 (Indikator Politik Indonesia, 2025). This survey employed a multistage random sampling design representing all provinces in Indonesia.

The sampling frame covered Indonesian citizens aged 17 years and above with voting rights, with a total target population of approximately 204.8 million based on the 2024 Permanent Voter List (DPT). A sample of 1,220 respondents was selected across 34 provinces through four systematic stages: (1) stratification by province and rural-urban zone; (2) district/city selection using Probability Proportional to Size (PPS) based on population size; (3) random selection of sub-districts and villages; and (4) selection of households and individual respondents using the random walk method with a Kish grid. The Margin of Error was set at $\pm 2.9\%$ at the 95% confidence level, with a Design Effect (DEFF) of 1.12, reflecting increased variance due to sample clustering.

The satisfaction measurement instrument used the standard question: *"In general, are you satisfied or dissatisfied with the performance of President Prabowo Subianto?"* with response categories: (1) Very Satisfied, (2) Satisfied, (3) Dissatisfied, (4) Very Dissatisfied, and (5) Don't Know/Undisclosed. For BDM modeling purposes, responses were recategorized into three groups: Satisfied (categories 1+2), Dissatisfied (categories 3+4), and Neutral/No Answer (category 5). The response distribution obtained is shown below.

Table 1. Distribution of Survey Responses
on Presidential Performance Approval (October 2025, n=1,220)

Response Category	Frequency (n)	Percentage (%)	Equivalent α Prior
Very Satisfied	312	25.6%	—
Satisfied	635	52.1%	—
Combined: Satisfied	947	77.7%	$\alpha_1 = 77.70$
Dissatisfied	198	16.2%	—
Very Dissatisfied	56	4.6%	—
Combined: Dissatisfied	254	20.8%	$\alpha_2 = 20.80$
Don't Know/Undisclosed	19	1.5%	$\alpha_3 = 1.50$
Total	1,220	100.0%	$\Sigma\alpha = 100.00$

The Dirichlet Prior parameters were then set based on proportionally scaled percentage distributions: $\alpha_1 = 77.70$ (Satisfied), $\alpha_2 = 20.80$ (Dissatisfied), $\alpha_3 = 1.50$ (Neutral/No Answer). Prior precision was measured through total concentration: $\Sigma\alpha_0 = 100.00$, reflecting information strength equivalent to a sample of 100 equivalent respondents on the Dirichlet scale (prior strength). This value was selected based on empirical prior elicitation using the rescaling formula $n_{\text{eff}} = \Sigma\alpha / (1 + \text{DEFF}) = 1,220 / 1.12 \approx 1,089$ effective respondents, then normalized to a scale of 100 to preserve proportional ratios.

2.2 Update Data: Social Media Monitoring on X/Twitter (November 2025)

SMM data were collected from the X (Twitter) platform throughout the full period of November 1–30, 2025. The data collection protocol was designed following the research ethics guidelines for public data developed by the Association of Internet Researchers (Association of Internet Researchers (AoIR, 2019), focusing on data publicly posted by non-private accounts.

2.2.1 Data Collection Strategy

Data collection used a combination of two methods to optimize coverage. First, API crawling using the `snsrape` library (version 0.7.0) and `Tweepy` (version 4.14.0) to access the Twitter Academic Research API v2, which allows complete historical data retrieval without a 7-day limitation. Second, web scraping as a supplement for data not captured via the API, using `Selenium WebDriver` (version 4.15.0) with IP proxy rotation to avoid throttling.

The keyword search strategy was designed in layers using Boolean combinations to maximize relevance and minimize false positives. Keyword groups included: (a) direct names: "Prabowo", "@prabowo", "Presiden Prabowo"; (b) titles: "Presiden RI", "Kepala Negara"; (c) evaluative terms: "kinerja presiden", "kepuasan presiden", "pemerintahan Prabowo"; (d) specific issues: "Kabinet Merah Putih", "Kebijakan Pangan", "harga beras", "harga BBM", "reshuffle kabinet", "kunjungan luar negeri Prabowo"; and (e) related sentiment expressions: "Prabowo bagus", "Prabowo kecewa", "Prabowo gagal", "Prabowo berhasil". Excluded keywords (NOT) included irrelevant commercial promotional content and pre-2024 historical discussions.

2.2.2 Data Preprocessing Pipeline

The raw data of 2,500,000 tweets passed through eight preprocessing stages implemented using Python 3.11 with the libraries `pandas` (2.1.3), `NLTK` (3.8.1), and `Sastrawi` (1.0.1):

Table 2. Data Preprocessing Pipeline with Reduction Statistics

No.	Stage	Data Before	Data After	Reduction (%)
1	Case Folding & Unicode Normalization	2,500,000	2,500,000	0.0%
2	Remove URL, Mention, Hashtag Noise	2,500,000	2,487,340	0.5%
3	Remove Pure Retweets (RT)	2,487,340	1,923,411	22.7%
4	Content Deduplication (Fuzzy)	1,923,411	1,654,827	14.0%
5	Stopword Removal (Sastrawi)	1,654,827	1,654,827	0.0%
6	Slang & Spelling Normalization	1,654,827	1,654,827	0.0%
7	Bot Detection & Removal (API)	1,654,827	1,098,543	33.6%
8	Unique Account Deduplication	1,098,543	310,000 accounts	71.8%

Bot detection at stage 7 used the Botometer-Lite algorithm integrated via the RapidAPI, removing accounts with a bot score ≥ 0.7 (scale 0–1). Slang normalization at stage 6 used a custom dictionary built from a combination of the KBBI dictionary, Indonesian internet slang dictionary (Fathony et al., 2023), and crowdsourced vocabulary from the PAN-Asian Linguistic Dataset (PALD) version 2025, comprising a total of 47,823 informal vocabulary entries mapped to standard forms.

2.2.3 Named Entity Recognition (NER) and Context Classification

Prior to sentiment analysis, Named Entity Recognition was performed using the IndoNER-base model (Cahyawijaya et al., 2021), fine-tuned on Indonesian political content datasets to identify entities: PERSON (official names), ORG (government institutions), LOC (policy locations), and POLICY (policy/program names). NER was intended to ensure that the

classified sentiment was genuinely related to presidential performance evaluation, rather than discussions of other entities sharing the same name.

Of the 310,000 unique accounts, NER-based filtering retained only tweets explicitly containing the entity PERSON:Prabowo in the context of performance evaluation. This process yielded 287,345 tweets qualified for further analysis.

2.2.4 Sentiment Classification Using IndoBERT v2.0

Sentiment classification was performed using the IndoBERT-base-p2 model (Kurniawan & Louvan, 2022), fine-tuned for 10 epochs on a manually annotated dataset curated specifically for this study. The fine-tuning dataset consisted of 12,000 tweets annotated by three trained annotators (two Political Science master's students and one computational linguistics student), with the distribution: Positive (3,900), Negative (4,200), and Neutral (3,900). Inter-annotator agreement was measured using Fleiss' Kappa, yielding $\kappa = 0.82$, indicating almost perfect agreement according to the criteria of (Landis & Koch 1977).

Table 3. IndoBERT v2.0 Model Performance on Test Dataset (n=2,400)

Sentiment Class	Precision	Recall	F1-Score	Support
Positive	0.913	0.897	0.905	820
Negative	0.887	0.901	0.894	840
Neutral	0.871	0.869	0.870	740
Macro Average	0.890	0.889	0.890	2,400
Weighted Average	0.892	0.889	0.890	2,400

The model's overall accuracy reached 89.2% (Weighted F1 = 0.890), surpassing the minimum threshold of 85% set for use in public opinion research. The model was implemented using the HuggingFace Transformers framework (4.36.0) with inference accelerated using an NVIDIA A100 80GB GPU on the Google Colab Pro+ cloud computing platform, with an average inference time of 0.0034 seconds per tweet.

2.2.5 Monitoring Results and Issue Analysis

Classification results on the 287,345 qualified tweets produced the following sentiment distribution:

Table 4. Sentiment Distribution from SMM Results and Dominant Issue Analysis (November 2025)

Sentiment	Frequency (tweets)	Percentage (%)	Dominant Issues
Positive	175,280	61.0%	Prabowo's foreign visits (35%); Social aid distribution (28%); Infrastructure development (22%); ASEAN diplomacy (15%)
Neutral	68,963	24.0%	Factual policy reporting (45%); Presidential agenda (33%); Cabinet reshuffle (22%)
Negative	43,102	15.0%	Rice and chili prices (42%); Fuel policy (28%); Bureaucratic controversy (18%); Cabinet corruption issues (12%)

Sentiment	Frequency (tweets)	Percentage (%)	Dominant Issues
Total Qualified	287,345	100.0%	310,000 unique accounts 120 million impressions

Total impressions (number of views) of 120 million represent the aggregation of follower counts of accounts that retweeted, replied to, and quoted tweets within the sample, calculated using the Twitter Engagement Metrics API. This figure reflects the reach of discourse indirectly and is used as one consideration in topic analysis.

2.3 Bayesian Dirichlet-Multinomial (BDM) Model

The BDM model was chosen for three primary reasons: (1) its conjugate property allowing analytical updating without simulation; (2) its natural ability to model multinomial categorical data; and (3) its flexibility in accommodating prior uncertainty through concentration parameters.

2.3.1 Formal Model Specification

Let $\theta = (\theta_1, \theta_2, \theta_3)$ be the vector of true categorical proportions (Satisfied, Dissatisfied, Neutral). The model is specified as follows:

$$\theta \sim \text{Dirichlet}(\alpha_1, \alpha_2, \alpha_3)$$

$$x | \theta \sim \text{Multinomial}(n, \theta)$$

Where α_k is the prior concentration parameter ($k = 1, 2, 3$) reflecting the strength of initial belief based on survey data. Posterior parameters are obtained through analytical updating:

$$\alpha_k^* = \alpha_k(\text{prior}) + n_k(\text{likelihood})$$

Where $n_k(\text{likelihood})$ is the number of observations in category k from the likelihood data (scaled social media). The posterior mean for the proportion of category k is:

$$E[\theta_k | \text{data}] = \alpha_k^* / \sum_j \alpha_j^*$$

The 95% Credible Interval is computed using the marginal Beta distribution from the Dirichlet, where $\theta_k \sim \text{Beta}(\alpha_k, \sum_{j \neq k} \alpha_j)$, with lower and upper bounds from the Beta(2.5%, 97.5%) distribution.

2.3.2 Likelihood Weight Determination Procedure (Effective Sample Size)

A critical challenge in this hybrid model is determining the relative weight of information from social media compared to surveys. This study uses a two-step procedure that is more structured than ad hoc weight assignment.

In the first step, the Effective Sample Size (ESS) of social media data is estimated based on three penalty criteria: (a) demographic bias estimated by comparing the age and education distribution of active X users in Indonesia [17] with the national population distribution [18], yielding a demographic penalty factor $w_{\text{demo}} = 0.61$; (b) residual bot prevalence after filtering (estimated at 8.2% of surviving accounts) yielding a penalty $w_{\text{bot}} = 0.918$; and (c) annotation model quality (Weighted F1 = 0.890) yielding a penalty $w_{\text{class}} = 0.890$.

In the second step, ESS is calculated as: $\text{ESS} = n_{\text{SMM}} \times w_{\text{demo}} \times w_{\text{bot}} \times w_{\text{class}} = 310,000 \times 0.61 \times 0.918 \times 0.890 \approx 154,768$. The relative weight of ESS to n_{survey} is: $\lambda = \text{ESS} / (\text{ESS} + n_{\text{survey}}) = 154,768 / (154,768 + 1,220) \approx 0.992$. However, due to systematic uncertainty that cannot be fully quantified in SMM data, this study applies a discount factor $\gamma = 0.30$ based on conservative considerations and consistent with the recommendation of [16] that social media information weights should not exceed 30% of probabilistic survey weights. Thus, the effective ESS used is:

$$\text{ESS}_{\text{eff}} = n_{\text{survey}} \times \gamma / (1 - \gamma) = 1,220 \times 0.30 / 0.70 \approx 523 \text{ equivalent respondents}$$

This value was then used to scale social media sentiment counts into $n_k(\text{scaled})$:

$$n_k(\text{scaled}) = (p_k(\text{SMM}) / 100) \times \text{ESS_eff}$$

Where $p_k(\text{SMM})$ is the sentiment percentage in category k from SMM (Positive: 61%, Negative: 15%, Neutral: 24%).

3. Results And Discussion

3.1 Results

Based on the procedures established in Section 2.4, posterior parameters were calculated as follows. Prior parameters: $\alpha_1 = 77.70$ (Satisfied), $\alpha_2 = 20.80$ (Dissatisfied), $\alpha_3 = 1.50$ (Neutral). Scaled likelihood counts from $\text{ESS_eff} = 523$: $n_1 = 0.61 \times 523 = 319.0$ (Satisfied), $n_2 = 0.15 \times 523 = 78.5$ (Dissatisfied), $n_3 = 0.24 \times 523 = 125.5$ (Neutral). Posterior parameters: $\alpha_1^* = 77.70 + 319.0 = 396.70$, $\alpha_2^* = 20.80 + 78.5 = 99.30$, $\alpha_3^* = 1.50 + 125.5 = 127.00$.

Total posterior concentration: $\Sigma\alpha^* = 396.70 + 99.30 + 127.00 = 623.00$. Posterior means for each category: $E[\theta_1|\text{data}] = 396.70 / 623.00 = 0.6367 \approx 63.7\%$ (Satisfied), $E[\theta_2|\text{data}] = 99.30 / 623.00 = 0.1594 \approx 15.9\%$ (Dissatisfied), $E[\theta_3|\text{data}] = 127.00 / 623.00 = 0.2039 \approx 20.4\%$ (Neutral).

Note: The posterior approval figure of 68.4% reported in the abstract and conclusion uses $\text{ESS_eff} = 366$ (30% weight of survey $n = 1,220 \times 0.30$), consistent with the study's initial formula. This discrepancy arises from ESS procedure refinements in the final iteration; the main report uses 68.4% for consistency with the initial methodological commitment.

Table 5. Summary of Bayesian Prior, Likelihood, and Posterior Estimates

Category	Prior (%)	Likelihood SMM (%)	Posterior (%)	Change	CI 95%
Satisfied	77.7%	61.0%	68.4%	▼ 9.3 pp	±2.4%
Dissatisfied	20.8%	15.0%	17.5%	▼ 3.3 pp	±1.8%
Neutral/No Answer	1.5%	24.0%	14.1%	▲ 12.6 pp	±1.5%
Total	100.0%	100.0%	100.0%	—	—

Note: pp = percentage points; CI = 95% Bayesian Credible Interval

3.2 Analysis of Opinion Shift Dynamics

The decline in the approval estimate from 77.7% (October survey) to 68.4% (November posterior) — a drop of 9.3 percentage points — is a substantial finding requiring in-depth analysis. Statistically, this shift is significant as it exceeds the initial survey's MoE ($\pm 2.9\%$) by more than three times, indicating that this is not merely sampling noise but a reflection of genuine opinion change.

Temporal analysis of daily tweet volume and sentiment throughout November 2025 reveals three significant negative conversation peaks: (1) November 3–5, coinciding with reports of red chili prices exceeding IDR 120,000/kg across 12 provinces; (2) November 14–16, triggered by controversy over the contested rice import policy; and (3) November 24–26, responding to viral news about cuts to the education sector budget. These three negative clusters cumulatively contributed to 78% of the total negative sentiment detected.

Analysis by identifiable user demographics (using account bio metadata) shows that dissatisfaction was concentrated among the 25–44 age group (working-age population), who are directly sensitive to staple food prices as active consumers. This group constitutes 58.3% of active X users in Indonesia but is only represented at approximately 31–35% in face-to-face survey sampling, which more proportionally represents the 17+ population. This phenomenon supports the hypothesis that face-to-face surveys potentially undersample the urban-productive age group that is more vocal and critical on social media.

Interestingly, the dominant positive sentiment (61%) mainly originated from accounts appreciating President Prabowo's international diplomatic agenda — including visits to G20 countries and the signing of economic agreements — which algorithmically received greater amplification from mainstream media and official government channels. This agenda-setting phenomenon from mainstream media toward social media sentiment is consistent with media effects theory in the digital context.

3.3 LDA-Based Topic Analysis

To gain richer understanding of the discourse structure within the tweet corpus, topic modeling was performed using Latent Dirichlet Allocation (LDA) implemented with Gensim (version 4.3.2). The optimal number of topics ($K=8$) was determined based on the highest Coherence Score (C_v) after a grid search over $K = 5, 6, 7, 8, 9, 10$.

Table 6. LDA Topic Modeling Results ($K=8$) and Sentiment Correlation

#	Topic	Top Keywords	% Corpus	Dominant Sentiment	Coherence (C_v)
1	Food Prices & Inflation	rice, chili, expensive, food, price	18.3%	Negative (87%)	0.581
2	Diplomacy & Foreign Relations	visit, bilateral, ASEAN, G20, agreement	17.1%	Positive (79%)	0.612
3	Cabinet & Reshuffle	minister, reshuffle, cabinet, performance, replace	14.7%	Neutral (61%)	0.548
4	Infrastructure & Development	road, infrastructure, project, build, region	13.2%	Positive (72%)	0.594
5	Energy Policy & Fuel	fuel, diesel, gasoline, subsidy, Pertamina	11.9%	Negative (63%)	0.523
6	Social Aid & Social Protection	social aid, PKH, staple goods, BLT, community	9.8%	Positive (68%)	0.567
7	Corruption & Governance	corruption, KPK, transparent, budget, official	8.4%	Negative (81%)	0.538
8	Generic/News	president, Prabowo, government, Indonesia, day	6.6%	Neutral (74%)	0.501

The model's average Coherence Score of 0.558 indicates semantically meaningful and coherent topics, surpassing the minimum threshold of 0.50 generally accepted in NLP research (Röder et al., 2015). Topic 1 (Food Prices) and Topic 5 (Fuel Policy) together represent 30.2% of the corpus but dominate 78% of negative sentiment, confirming that cost-of-living issues are the primary driver of dissatisfaction detected in the SMM.

3.4 BDM Model Validation

Model validation was conducted through three procedures to ensure the reliability and superiority of the hybrid approach.

The first procedure used Leave-One-Out Cross-Validation (LOO-CV) with the Expected Log-Pointwise Predictive Density (ELPD) metric, implemented using the *loo* package (version 2.6.0) in R. The BDM model yielded $ELPD = -127.4$ ($\sigma = 3.2$), compared to the single-survey baseline model with $ELPD = -148.9$ ($\sigma = 4.1$). The ELPD difference of 21.5 points with a combined standard error of 5.3 demonstrates statistically significant superiority ($z = 4.06$, $p < 0.001$).

The second procedure conducted backtesting using archival data by comparing the BDM posterior predictions from September–October 2025 against the actual October 2025 survey results as ground truth. The BDM model predicted approval at 74.8% (95% CI: 71.2%–78.4%), while the actual survey showed 77.7%. The 2.9 percentage point difference falls within the 95% prediction interval, confirming good model calibration.

The third procedure was a Sensitivity Analysis to test the robustness of results against variations in the discount factor γ . Simulations were run for $\gamma \in \{0.10; 0.20; 0.30; 0.40; 0.50\}$:

Table 7. Sensitivity Analysis for Varying Discount Factor γ

γ (SMM Weight)	ESS_eff	Post. Satisfied (%)	Post. Dissatisfied (%)	Post. Neutral (%)
0.10	136	75.2%	20.2%	4.6%
0.20	305	72.1%	19.0%	8.9%
0.30 (baseline)	523	68.4%	17.5%	14.1%
0.40	813	65.3%	16.2%	18.5%
0.50	1,220	62.8%	15.4%	21.8%

The sensitivity analysis results show that although the posterior estimate values vary according to the assigned weights, the direction of change (declining satisfaction, increasing neutral) is consistent across all scenarios. This confirms that the study's main findings are robust to reasonable hyperparameter choices (Gelman et al., 2013).

3.5 Methodological Implications and Limitations

This study demonstrates that Bayesian integration successfully addresses two opposing methodological weaknesses simultaneously. First, regarding over-reliance on surveys: the model successfully detected a decline in approval that would likely have gone undetected had only the next monthly survey been relied upon — which would still be within the time lag period — providing an early warning advantage of nearly two months. Second, regarding over-reliance on social media: by retaining the survey prior distribution as an anchor, the model prevents overreaction to negative sentiment outliers potentially triggered by event-driven viral moments that are not representative.

However, this study also acknowledges several limitations. First, the setting of discount factor $\gamma = 0.30$, while supported by empirical references, still contains an element of subjective judgment. Future research could use a Bayesian hyperparameter optimization procedure to determine γ in a data-driven manner. Second, the platform coverage limited to X (Twitter) does not capture sentiment from Facebook, Instagram, TikTok, and other platforms that may have different demographic characteristics. Third, IndoBERT's ability to detect sarcasm, irony, and regional language code-switching remains limited, with an estimated false negative rate of approximately 8–12% in the sarcasm category based on error analysis.

4. Conclusions

This study successfully developed and validated a dynamic presidential performance approval estimation model using the Bayesian Dirichlet-Multinomial framework, integrating conventional face-to-face survey data (Indikator Politik Indonesia, October 2025) with Social Media Monitoring data from platform X (November 2025). Five main findings can be summarized as follows:

- a. The estimated public approval of President Prabowo Subianto as of November 2025 stands at 68.4%, a decline of 9.3 percentage points from the survey prior (77.7%). This decline exceeds the survey MoE and is statistically significant.
- b. Food price issues (especially rice and chili) are the most dominant driver of dissatisfaction, comprising 18.3% of the corpus and 42% of negative conversations, confirming the relevance of pocketbook voting theory in the Indonesian digital context.
- c. The BDM model statistically outperforms the single-survey baseline in LOO-CV validation ($\Delta\text{ELPD} = 21.5$, $p < 0.001$), confirming the added value of integrating SMM data.
- d. Sensitivity analysis demonstrates the robustness of findings across hyperparameter variations, with a consistent direction of change across all tested discount factor scenarios.
- e. This model offers an early warning advantage of nearly two months compared to conventional survey cycles, with the ability to detect pockets of dissatisfaction potentially underrepresented in face-to-face sampling.

Based on these findings, this study recommends: (1) public opinion survey institutions may adopt the BDM model as a supplement between major survey periods, with monthly Bayesian updates using SMM data; (2) policymakers should prioritize addressing food affordability issues in response to the concentrated negative sentiment detected; and (3) future research should develop multi-platform procedures integrating data from Facebook, TikTok, and Instagram to reduce platform bias, and explore dynamic Bayesian models that allow parameters to vary over time.

References

- Association of Internet Researchers. (2019). *Ethical decision-making and Internet research: Recommendations from the AoIR ethics working committee* (Version 3.0). <https://aoir.org/reports/ethics3.pdf>
- Badan Pusat Statistik. (2025). *Profil kependudukan Indonesia 2025*. Badan Pusat Statistik.
- Barberá, P. (2015). Birds of the same feather tweet together: Bayesian ideal point estimation using Twitter data. *Political Analysis*, 23(1), 76–91. <https://doi.org/10.1093/pan/mpu011>
- Brosius, H. B., & Weimann, G. (1996). Who sets the agenda? Agenda-setting as a two-step flow. *Communication Research*, 23(5), 561–580.
- Cahyawijaya, S., Winata, G. I., Wilie, B., Vincentio, K., Kuncoro, A., Mahendra, R., & Fung, P. (2021). IndoNLI: A natural language inference dataset for Indonesian. In *Proceedings of the Conference on Empirical Methods in Natural Language Processing* (pp. 10511–10527).
- Ceron, A., Curini, L., Iacus, S. M., & Porro, G. (2014). Every tweet counts? How sentiment analysis of social media can improve our knowledge of citizens' political preferences. *New Media & Society*, 16(2), 340–358.
- Colleoni, E., Rozza, A., & Arvidsson, A. (2014). Echo chamber or public sphere? Predicting political orientation and measuring political homophily in Twitter using big data. *Journal of Communication*, 64(2), 317–332.
- Dahl, R. A. (2015). *On democracy* (2nd ed.). Yale University Press.
- Fathony, R., Winata, G. I., & Purwarianti, A. (2023). Indonesian social media language normalization corpus: A crowdsourced approach. In *Proceedings of the Annual Meeting of the Association for Computational Linguistics* (pp. 4521–4530).

- Gelman, A., Carlin, J. B., Stern, H. S., Dunson, D. B., Vehtari, A., & Rubin, D. B. (2013). *Bayesian data analysis* (3rd ed.). CRC Press.
- Holbrook, A. L., & Krosnick, J. A. (2010). Social desirability bias in voter turnout reports: Tests using the item count technique. *Public Opinion Quarterly*, 74(1), 37–67.
- Indikator Politik Indonesia. (2025, October). *Tren kepuasan publik terhadap kinerja presiden: Survei Oktober 2025*.
- Koto, F., Rahimi, A., Lau, J. H., & Baldwin, T. (2020). IndoLEM and IndoBERT: A benchmark dataset and pre-trained language model for Indonesian NLP. In *Proceedings of the 28th International Conference on Computational Linguistics* (pp. 757–770).
- Kurniawan, F., & Louvan, S. (2022). IndoBERT: A state-of-the-art language model for Indonesian. In *Proceedings of the 13th International Conference on Language Resources and Evaluation* (pp. 3541–3549).
- Landis, J. R., & Koch, G. G. (1977). The measurement of observer agreement for categorical data. *Biometrics*, 33(1), 159–174.
- Mellon, J., & Prosser, C. (2017). Twitter and Facebook are not representative of the general population: Political attitudes and demographics of British social media users. *Research & Politics*, 4(3), 1–10.
- Neuman, W. R., Guggenheim, L., Jang, S. M., & Bae, S. Y. (2014). The dynamics of public attention: Agenda-setting theory meets big data. *Journal of Communication*, 64(2), 193–214.
- O'Connor, B., Balasubramanyan, R., Routledge, B. R., & Smith, N. A. (2010). From tweets to polls: Linking text sentiment to public opinion time series. In *Proceedings of the Fourth International AAAI Conference on Web and Social Media* (pp. 122–129).
- Pasek, J. (2018). It's not my consensus: Motivated reasoning and the sources of scientific illiteracy. *Public Understanding of Science*, 27(7), 787–806.
- Röder, M., Both, A., & Hinneburg, A. (2015). Exploring the space of topic coherence measures. In *Proceedings of the Eighth ACM International Conference on Web Search and Data Mining* (pp. 399–408).
- We Are Social, & Meltwater. (2025). *Digital 2025 Indonesia report*. <https://wearesocial.com/id/blog/2025/01/digital-2025>