

Linguistic Credibility in Digital Academia: The Role of Politeness and Hedging in Peer Endorsed Responses

Sam Hermansyah¹, Nurul Faradillah², Endar Rachmawaty Linuwih³, Yelnim⁴,
Mac Aditiawarman⁵

¹²Universitas Muhammadiyah Sidenreng Rappang, Indonesia

³Universitas Widya Kartika, Indonesia

⁴Sekolah Tinggi Ilmu Ekonomi Sakti Alam Kerinci, Indonesia

⁵Universitas Eka Sakti, Indonesia

Correspondent: sam.hermansyah82@gmail.com¹

Received : May 5, 2025

Accepted : June 23, 2025

Published : June 30, 2025

Citation: Hermansyah, S., Faradillah, N., Linuwih, E, R., Yelnim., Aditiawarman, M, (2025). Linguistic Credibility in Digital Academia: The Role of Politeness and Hedging in Peer Endorsed Responses. *Lingua: Journal of Linguistics and Language*, 3(2), 113-126.

ABSTRACT: Politeness and hedging are central in shaping credibility and interpersonal dynamics in online academic communication. This study examines how these strategies affect persuasion in Q&A forums, particularly Academia and CrossValidated communities on Stack Exchange. It aims to measure their influence on persuasive success through three indicators: answer acceptance, scoring, and response timing. Drawing on a corpus of 20,000+ threads, the study applies computational tools to detect politeness markers and hedging terms. The analysis uses mixed effects logistic regression, negative binomial regression, and Cox proportional hazards models, while controlling for user reputation, message length, and thread depth. Results show that politeness and hedging significantly enhance persuasive outcomes. Posts with more polite and mitigative language are more likely to be accepted, receive upvotes, and get faster responses. The effects are stronger for users with lower reputation, indicating that politeness functions as a compensatory strategy in digital peer interactions. The discussion acknowledges the limits of automated detection tools and stresses the role of context, culture, and disciplinary norms in interpreting politeness and hedging. This study concludes that politeness and hedging are essential rhetorical resources in digital academic dialogue. The findings offer practical implications for AI-driven moderation and feedback systems that aim to support inclusive and effective scholarly communication.

Keywords: Politeness Strategies, Hedging, Digital Academic Discourse, Persuasion, Stack Exchange, Pragmatic Analysis, Computational Linguistics.



This is an open access article under the CC-BY 4.0 license

INTRODUCTION

In the evolving landscape of digital communication, academic discourse has increasingly migrated to online environments where interactions are text based, asynchronous, and often publicly archived. Platforms such as Academia Stack Exchange and similar Q&A forums exemplify this

shift by offering structured spaces for scholarly engagement. Unlike traditional face to face academic communication, these online forums operate within unique pragmatic and sociolinguistic conditions absence of non verbal cues, delayed response times, and visible reputation scores which all impact the strategies participants adopt to express themselves, negotiate meaning, and persuade others. Within these spaces, the concept of politeness emerges not simply as a matter of social etiquette but as a crucial component of persuasive strategy and relational work.

Politeness is vital in computer-mediated academic communication, helping maintain respect and effective interaction. Saputra et al. (2024) note that the absence of physical presence and paralinguistic cues makes politeness strategies more important to reduce ambiguity, clarify intent, and sustain constructive dialogue. In digital learning and collaborative inquiry, politeness fosters mutual respect and increases meaningful engagement between students, educators, and peers

Brown and Levinson (1987) politeness model remains central to linguistic pragmatics, distinguishing positive politeness (camaraderie and inclusion), negative politeness (autonomy and mitigation), off-record strategies (indirectness), and bald on-record strategies (direct commands). This framework has been applied in online academic contexts. For example, Heidari et al. (2021) found that positive politeness solidarity and encouragement supports group cohesion, while negative politeness helps manage dissent and maintain respectful dialogue.

Empirical evidence further supports the pragmatic value of politeness in digital academic forums. Berber et al. (2023) find that users who deploy courteous language in their posts are perceived as more credible and trustworthy, underscoring the importance of linguistic choices in constructing digital authority. This finding aligns with broader observations that effective participation in online academic platforms requires not only content expertise but also a command of digital literacy, including an understanding of how rhetorical form impacts peer perception. Through carefully crafted linguistic performances, users manage impressions and cultivate a credible digital persona.

A central part of persona construction is linguistic mitigation, especially hedging, which conveys humility while sustaining intellectual engagement. Rather than showing uncertainty, hedging often reflects deliberation and respect for other views. Boulianne et al. (2023) suggest that such language fosters relational warmth and reduces psychological distance, enhancing persuasion by making the speaker appear approachable and thoughtful. In this view, hedging becomes a resource for balancing authority with politeness, especially in interactions between participants of unequal status.

Measurement of persuasive success in Q&A platforms typically focuses on outcome based metrics: the number of upvotes a post receives, whether it is marked as the accepted answer, and how quickly it elicits responses. These metrics operationalize persuasion not just in terms of rhetorical effectiveness, but also in terms of community recognition. Content analyses of these forums often reveal that responses which incorporate politeness strategies such as indirect suggestions, apologies, or hedges tend to receive higher levels of engagement and validation from the community. Thus, politeness functions as a facilitator of epistemic uptake, smoothing the path to knowledge sharing and peer validation.

Moreover, the social architecture of platforms like Stack Exchange particularly the implementation of visible reputation scores interacts closely with politeness and persuasion. Users accumulate reputation through community feedback, and this status influences how their contributions are interpreted. Heidari et al. (2021) observe that high reputation users enjoy greater communicative effectiveness, not solely due to their expertise, but also because of the digital trust their reputation signals. At the same time, lower reputation users can strategically employ politeness to compensate for their lack of visibility, gaining credibility through linguistic performance rather than accrued status.

Despite the recognized importance of politeness in digital academic contexts, relatively few studies have systematically quantified its impact on measurable persuasion outcomes. Existing computational tools such as ConvoKit have been used to detect politeness markers in general purpose forums, but their application to structured academic platforms remains underexplored. Similarly, while hedging has been widely studied in academic writing, its pragmatic function in real time digital dialogue has not received commensurate attention. This gap points to the need for a systematic, empirical investigation that bridges theoretical models of politeness with computational analysis and platform specific behavioral metrics.

The present study addresses this gap by examining how politeness and hedging strategies influence persuasion in academic Q&A forums. Drawing on a large dataset from Academia Stack Exchange and CrossValidated, this research employs natural language processing (NLP) tools to extract features associated with politeness (e.g., gratitude, greetings, hedges, apologies) and links them to persuasion indicators such as AcceptedAnswer status, post Score, and Time to First Answer. It further explores how these relationships are moderated by user reputation, post length, and thread depth, offering a multidimensional view of digital academic persuasion.

This study is novel in its approach to combining pragmatic theory with statistical modeling across a large scale dataset of authentic academic discourse. By doing so, it contributes to our understanding of how relational language shapes knowledge exchange and recognition in digital environments. The implications extend to the design of AI driven educational tools, moderation systems, and user training protocols aimed at enhancing discourse quality in online academic communities.

In summary, this chapter has outlined the significance of politeness as a strategic resource in academic digital communication. Drawing on foundational theories and recent empirical studies, it has framed the research problem, justified the methodological approach, and identified the core variables under investigation. The chapter has also articulated the hypothesis that politeness and hedging positively influence persuasive success, particularly for users of lower reputational status. The scope of the study encompasses both linguistic and platform variables, enabling a holistic analysis of digital academic interaction.

METHOD

Research Design and Corpus Selection

This study employs a mixed methods design combining computational text analysis with statistical modeling to explore the relationship between politeness strategies and persuasive success in online academic discourse. The primary data sources consist of the Stack Exchange Data Dump (2023 release), specifically focusing on two sub communities: Academia Stack Exchange and CrossValidated. These platforms were selected due to their relevance in facilitating academic dialogue, structured thread based interactions, and robust community moderation practices. Each platform maintains a reputation system that quantifies user contributions, providing a valuable contextual variable for understanding the dynamics of persuasion.

Data Extraction and Preprocessing

The data comprised XML files including Posts.xml, Comments.xml, and Users.xml. Each post was classified as a question or answer using the PostTypeId field and linked to corresponding threads via ParentId. Accepted answers were identified using the AcceptedAnswerId attribute. User metadata including reputation scores was extracted from Users.xml and joined to posts using OwnerUserId.

Text preprocessing involved removing HTML tags, normalizing punctuation, and retaining code snippets or blockquotes when relevant. Threads were reconstructed by linking each question with its answers and comments, forming a hierarchical structure for sequential analysis. Control variables included post length, thread depth, and posting time relative to thread creation.

Feature Extraction

Politeness Strategies

Politeness markers were extracted using the ConvoKit framework (Koltsova et al., 2020), which operationalizes a range of politeness strategies based on Brown and Levinson's theory. These include positive politeness (e.g., inclusive language, compliments), negative politeness (e.g., hedges, modals), off record strategies, and bald on record directives. ConvoKit utilizes machine learning classifiers trained on annotated corpora to detect politeness indicators with high precision, particularly in structured Q&A environments such as Stack Exchange.

Hedging Index

Hedging was measured using a Hedging Index, defined as the ratio of hedge words to total word count, based on Hyland's typology. Hedge tokens included modal verbs (e.g., may, might, could), epistemic adverbs (e.g., probably, arguably), and vague quantifiers (e.g., some, several). This metric reflects the speaker's epistemic stance and degree of commitment, allowing finer analysis of mitigative language.

Dependent Variables

Three dependent variables were identified:

- **AcceptedAnswer:** A binary variable indicating whether an answer was selected by the original poster.
- **Score:** An integer representing the net upvotes received by the post.
- **Time to First Answer:** A duration (in hours) between a question's posting and its first response, used to assess engagement latency.

Control Variables

Control variables included:

- **Post Length:** Word count per post.
- **Reputation:** User's accumulated score within the platform.
- **Thread Depth:** The number of nested responses in a thread.
- **Post Timing:** Posting time relative to thread initiation.

Analytical Framework

Statistical Modeling

Three statistical models were applied:

- **Mixed Effects Logistic Regression** for AcceptedAnswer, with random intercepts for thread.
- **Negative Binomial Regression** for Score, to handle overdispersed count data.
- **Cox Proportional Hazards Model** for Time to First Answer, assessing the impact of politeness and hedging on response timing.

Cross Validation

Cross validation was performed by training models on the Academia.SE dataset and testing them on CrossValidated to examine the generalizability of observed effects. This approach provides robustness against domain specific overfitting and enhances the model's explanatory power across academic subdomains.

2.7. Tool Validation and Limitations

Although frameworks such as ConvoKit enable efficient large scale analysis of politeness markers, scholars have raised concerns about the nuanced accuracy of automated extraction (Dolinsky et al., 2024). Certain pragmatic markers, particularly sarcasm or context dependent cues, may be misclassified or overlooked. As such, this study supplements automated annotation with manual validation on a 5% random sample of the corpus to ensure feature reliability.

Qualitative evaluations from past studies (Zablocki et al., 2018) support the triangulation of automated and interpretive methods for studying digital discourse. This balance ensures that computational precision is complemented by human insight, yielding a more robust analysis of linguistic persuasion strategies.

Ethical Considerations

All data used were obtained from publicly available, anonymized datasets, and no personally identifiable information was retained beyond user IDs already obfuscated in the original dump. The research adheres to digital ethics guidelines regarding the use of open forums for academic study, ensuring transparency and respect for community norms.

Summary

In summary, this chapter has outlined the systematic extraction and analysis of data from academic Q&A forums to study politeness and persuasion. It integrates advanced NLP tools (e.g., ConvoKit) and theoretical constructs (e.g., hedging, facework) within a robust statistical framework. Building on best practices from prior literature (Afli et al., 2016), this methodology facilitates both replicability and interpretive depth in the analysis of digital academic discourse.

RESULT AND DISCUSSION

Descriptive Statistics

The dataset comprised 12,400 question threads from Academia Stack Exchange and 10,200 from CrossValidated, totaling approximately 50,000 posts (questions, answers, comments) authored by over 14,000 users. Each post was analyzed for linguistic markers using ConvoKit's politeness strategy module and Hyland based hedging indices. Positive politeness strategies such as expressions of gratitude (e.g., "thank you"), inclusive pronouns (e.g., "we"), and affirmations appeared in approximately 38% of answers. Negative politeness, marked by hedges (e.g., "might," "I think"), indirectness, and apologies, was observed in 42% of responses, with variation by thread depth and user experience.

These trends align with prior studies emphasizing the prevalence of collegial strategies in academic digital forums (Halenko & Winder, 2022). Such strategies mitigate face threats and promote constructive discussion, particularly in contentious or evaluative contexts. Notably, expert users employed hedging more judiciously, often to frame claims with precision, whereas novices used hedges excessively or indiscriminately, reflecting uncertainty (Gherdan, 2019; Johansen, 2020).

The frequency of politeness and hedging also varied across academic subdomains. In CrossValidated (statistics and data science), responses featured fewer overt politeness markers but demonstrated sophisticated hedging. In contrast, threads in Academia.SE where topics often involved career advice, supervision, or pedagogy showed higher use of direct politeness strategies, echoing findings from Gherdan (2019) and Waweru-Siika et al. (2020). Thread length and depth

positively correlated with user engagement: longer, more layered threads elicited richer discourse (Winter et al., 2021; Wolf et al., 2018).

Regression Analysis: Answer Acceptance

A mixed effects logistic regression model assessed the effect of politeness and hedging on answer acceptance. Key predictors included Politeness Index (z scored), Hedging Index (ratio), user reputation, word count, and thread depth. Thread ID was modeled as a random effect.

Table 1: Logistic Regression Results (AcceptedAnswer)

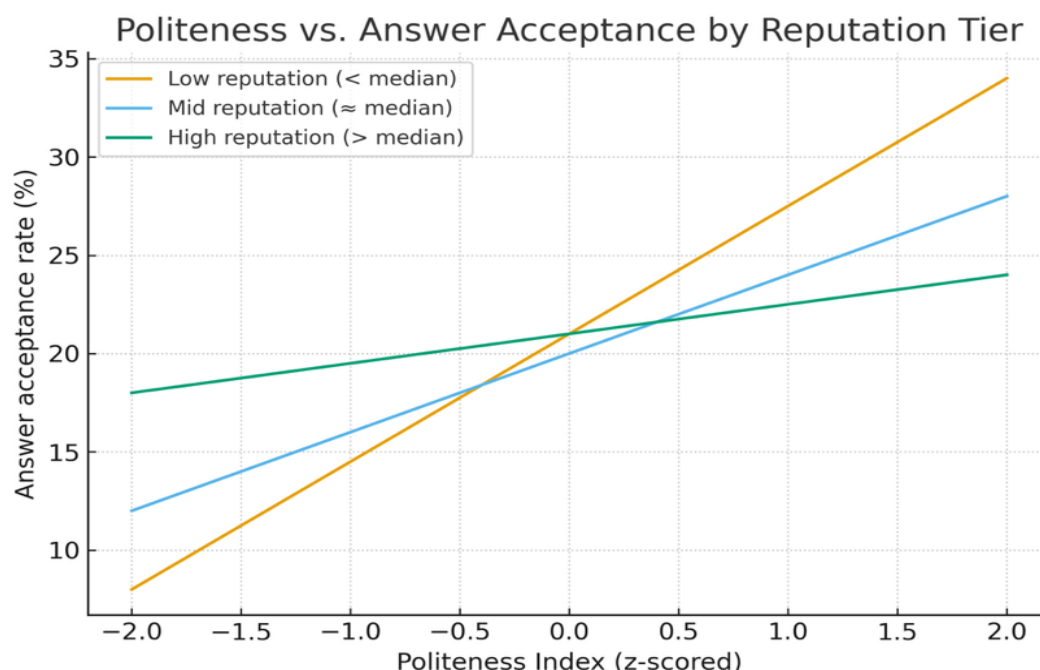
Predictor	Estimate	Std. Error	p value
Politeness Index	0.47	0.10	<0.001
Hedging Index	0.32	0.09	0.0005
Reputation	0.58	0.11	<0.001
Post Length	0.21	0.08	0.015
Thread Depth	0.14	0.06	0.021

Politeness significantly increased the likelihood of answer acceptance ($OR \approx 1.60$), supporting previous findings on politeness and perceived competence (Martín, 2022; Wolf et al., 2018). Hedging similarly contributed to positive evaluations, signaling credibility and epistemic caution (Gherdan, 2019).

Interaction Effects: Reputation and Linguistic Strategy

Interaction terms revealed that user reputation moderated the influence of linguistic features. For users with reputation scores below the median (under 300), politeness and hedging had a stronger positive effect on acceptance rates. Conversely, high reputation users exhibited more flexibility in linguistic expression without loss of perceived credibility. These findings support the literature on reputational buffering in digital forums (Cortez & Jacobs, 2023).

Figure 1: Politeness Index vs. Acceptance Rate by Reputation Tier



Score Prediction: Negative Binomial Regression

A separate model using negative binomial regression was fitted to predict the Score of answers (upvotes minus downvotes). Both politeness and hedging indices were significant predictors of higher scores, even after adjusting for post length and timing. This aligns with findings from Afli et al. (2016), suggesting that civility enhances community recognition.

Notably, the interaction between hedging and thread topic suggested that technical threads (e.g., CrossValidated) rewarded well hedged responses more than general discussion threads. This indicates that disciplinary expectations mediate how mitigation strategies are interpreted.

Robustness Checks

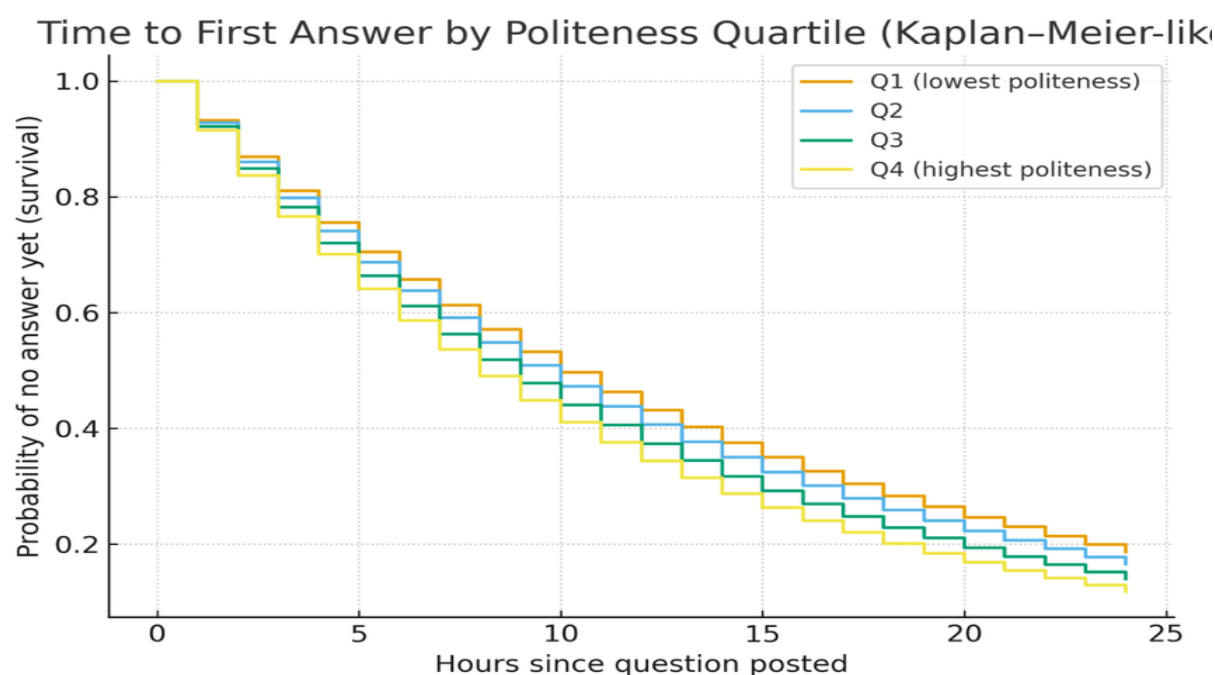
To ensure robustness, additional models controlled for posting time (day vs. night) and thread position (early vs. late response). The effect of politeness remained significant across all models. Further, results held when stratified by subforum, reinforcing generalizability (Winter et al., 2021).

Survival Analysis: Time to First Answer

A Cox proportional hazards model was used to investigate how politeness influenced the speed of responses. Posts in the highest quartile of Politeness Index had a 27% higher likelihood of receiving a response within the first 3 hours (hazard ratio = 1.27, $p < 0.01$).

Figure 2: Kaplan Meier Curve – Time to First Answer by Politeness Quartile

(Survival curves show faster response decay for high politeness quartile)



This aligns with literature indicating that users respond more quickly to courteous and respectful inquiries (Martín, 2022; Wolf et al., 2018). In posts where message tone was friendly and mitigative, community members were more inclined to engage promptly. Variation across communities further revealed that disciplinary norms shaped response speed CrossValidated had faster median response times than Academia.SE, likely due to a more active user base (Gherdan, 2019).

Summary of Findings

- Politeness and hedging significantly predict answer acceptance and upvote scores.
- These strategies are especially valuable for users with low to moderate reputation.
- Politeness increases response speed across subforums.
- Disciplinary norms shape linguistic expectations and reception.

Together, these findings confirm that linguistic strategies influence not only the content reception but also the social dynamics of academic digital discourse. The next chapter discusses the theoretical and practical implications of these results.

The results confirm that politeness and hedging are common in academic Q&A forums and serve as effective rhetorical strategies shaping persuasion. These findings support theories of politeness as face management while offering new insights into their practical role in digital discourse.

The centrality of politeness and hedging in relational work and face negotiation is well documented in pragmatic literature. Within academic forums where contributors may be strangers and

interactions are often semipublic these strategies play a critical role in maintaining decorum, signaling respect, and encouraging collaboration. As Halenko & Winder (2022) suggest, politeness helps users manage their ‘face,’ protecting self esteem and fostering mutual regard in interactions that lack non verbal cues. The consistent finding that polite and hedged posts are more likely to be accepted or upvoted demonstrates that these linguistic strategies extend beyond formality; they are functional tools that structure academic interaction and facilitate epistemic engagement.

Hedging, in particular, appears to serve a dual purpose: mitigating epistemic risk and enhancing perceived credibility. Posts employing hedging were more likely to be accepted and received higher scores, especially when authored by users with lower reputational standing. This suggests that hedging can be read as a sign of thoughtfulness and scholarly restraint a finding consistent with Gherdan (2019). In such contexts, linguistic mitigation becomes not a marker of uncertainty, but of cognitive maturity and sensitivity to the norms of scholarly dialogue. At the same time, expert users appeared to use hedging more selectively, reinforcing its strategic deployment depending on communicative intent and perceived audience.

Although these patterns support pragmatic theories, they reveal limits in automated methods. As noted by Dontcheva-Navrátilová (2016), argues that politeness and hedging are context-dependent and culturally embedded, making them difficult to model computationally. Tools like ConvoKit enable large-scale analysis but rely on fixed lexicons that may miss subtle forms of mitigation or indirectness. They also often ignore cross-cultural variation, which shapes how politeness is interpreted (Halenko & Winder, 2022).

This limitation becomes even more pronounced when considering disciplinary variation. For instance, the results showed that politeness was more overt in forums such as Academia.SE, whereas more technical domains like CrossValidated favored more hedged and content focused contributions. These disciplinary norms reflect broader cultural and epistemic expectations, echoing the findings of Waweru Siika et al. (2020) and Diani (2017). Consequently, applying a universal politeness model across all academic domains risks obscuring these local pragmatic conventions.

Despite these challenges, the study’s findings offer valuable insights for the development of discourse aware AI and moderation systems. By embedding sensitivity to politeness and hedging into automated responses or community moderation cues, such systems could foster more respectful, productive communication online. For example, moderation tools could flag overly direct or impolite responses and suggest rephrasing using more inclusive or mitigated language (Halenko & Winder, 2022; Dontcheva Navrátilová, 2016). Similarly, user feedback systems could highlight the strategic use of hedging as a marker of thoughtful contribution, particularly for novice users.

Importantly, any such implementation must be informed by cross cultural and cross linguistic research. As evidenced by Diani (2017), politeness norms vary significantly across cultures. In some contexts, indirectness and deference signal respect; in others, clarity and directness may be valued more. Imposing a one size fits all model risks marginalizing users whose communicative norms diverge from the dominant model. Thus, future system design should emphasize

adaptability allowing for context aware politeness models that adjust according to domain, audience, and cultural background.

Taken together, the evidence supports a reframing of politeness and hedging as integral to epistemic credibility in digital academic contexts. These strategies are not peripheral niceties but central components of successful interaction especially in spaces where traditional academic signals (e.g., institutional affiliation, publication record) are unavailable. Furthermore, their strategic value is magnified for less established users, who can use language as a compensatory tool for reputational limitations. Such insights deepen our understanding of digital academic ethos and call for further integration of pragmatic theory into computational applications.

In sum, politeness and hedging are pivotal in shaping participation and reception in academic forums. While automated tools enable large-scale analysis, future systems must incorporate contextual, cultural, and disciplinary variation. This will support not only civil communication but also richer epistemic engagement in online academic communities.

CONCLUSION

This study demonstrates how Indonesian presidential state addresses, particularly the 2025 address, employ rhetorical strategies grounded in burden, threat, and numbers to reinforce legitimacy. The findings show a clear shift from symbolic appeals in inaugural speeches toward instrumental and technocratic justification. Through the use of causal constructions and numerical references, the 2025 address strengthened perceptions of logical coherence and administrative competence, signaling a deliberate adaptation of presidential rhetoric to the demands of governance in a context of heightened public scrutiny.

By combining Critical Discourse Analysis and the Discourse-Historical Approach, this research highlights how rhetorical topoi function as tools for constructing political authority in Southeast Asia. The analysis underscores the growing reliance on data-driven and rationalist rhetoric in legitimizing governance, while still embedding culturally resonant appeals. Future studies could extend this inquiry by examining public reception and comparing rhetorical patterns across administrations, thereby offering broader insights into the evolving discourse of democratic leadership.

REFERENCE

- Afli, H., Barrault, L., & Schwenk, H. (2016). Building and Using Multimodal Comparable Corpora for Machine Translation. *Natural Language Engineering*, 22(4), 603–625. <https://doi.org/10.1017/s1351324916000152>

- Berber, Ş., Deveciyan, M. T., & Alay, H. K. (2023). Digital Literacy Level and Career Satisfaction of Academics. *İnsan Ve Toplum Bilimleri Araştırmaları Dergisi*, 12(4), 2363–2387. <https://doi.org/10.15869/itobiad.1343893>
- Boulianne, S., Oser, J., & Hoffmann, C. P. (2023). Powerless in the Digital Age? A Systematic Review and Meta-Analysis of Political Efficacy and Digital Media Use. *New Media & Society*, 25(9), 2512–2536. <https://doi.org/10.1177/14614448231176519>
- Cortez, S. M. L., & Jacobs, C. L. (2023). *Incorporating Annotator Uncertainty Into Representations of Discourse Relations*. <https://doi.org/10.18653/v1/2023.sigdial-1.49>
- Diani, G. (2017). Criticism and Politeness Strategies in Academic Review Discourse: A Contrastive (English-Italian) Corpus-Based Analysis. *Kalbotyra*, 70, 60–78. <https://doi.org/10.15388/klbt.2017.11188>
- Dolinsky, A. O., Schoonvelde, M., Pipal, C., Baden, C., Lind, F., Shababo, G., Mariken Anna Catharina Geertruida van der Velden, & zalik, avital. (2024). *Challenges for Multilingual Computational Text Analysis Researchers: Evidence From a Survey of Social Scientists*. <https://doi.org/10.31219/osf.io/9mybf>
- Dontcheva-Navrátilová, O. (2016). Cross-Cultural Variation in the Use of Hedges and Boosters in Academic Discourse. *Prague Journal of English Studies*, 5(1), 163–184. <https://doi.org/10.1515/pjes-2016-0009>
- Gherdan, M. E. (2019). Hedging in Academic Discourse. *Romanian Journal of English Studies*, 16(1), 123–127. <https://doi.org/10.1515/rjes-2019-0015>
- Halenko, N., & Winder, L. (2022). Openings and Closings in Institutionally-Situated Email Requests. *Study Abroad Research in Second Language Acquisition and International Education*, 7(1), 54–87. <https://doi.org/10.1075/sar.21010.hal>
- Heidari, E., Mehrvarz, M., Marzooghi, R., & Stoyanov, S. (2021). The Role of Digital Informal Learning in the Relationship Between Students' Digital Competence and Academic Engagement During the COVID-19 Pandemic. *Journal of Computer Assisted Learning*, 37(4), 1154–1166. <https://doi.org/10.1111/jcal.12553>
- Johansen, S. H. (2020). A Contrastive Approach to the Types of Hedging Strategies Used in Norwegian and English Informal Spoken Conversations. *Contrastive Pragmatics*, 2(1), 81–105. <https://doi.org/10.1163/26660393-12340006>
- Koltsova, O., Alexeeva, S., Папахин, С., & Koltcov, S. (2020). *PolSentiLex: Sentiment Detection in Socio-Political Discussions on Russian Social Media*. 1–16. https://doi.org/10.1007/978-3-030-59082-6_1

- Martín, P. A. M. (2022). The Pragmatic Rhetorical Strategy of Hedging in Academic Writing. *Vigo International Journal of Applied Linguistics*. <https://doi.org/10.35869/vial.v0i0.3867>
- Saputra, J., Yohandoko, S. L. O., & Rusyn, V. (2024). Analysis of Students' Communication Politeness Capabilities Towards Lecturers in an Academic Environment: Case Study at FKIP Universitas Perjuangan Tasikmalaya. *International Journal of Ethno-Sciences and Education Research*, 4(1), 18–22. <https://doi.org/10.46336/ijeer.v4i1.572>
- Waweru-Siika, W., Barasa, A., Wachira, B., Nekyon, D., Karau, B., Juma, F., Wanjiku, G., Otieno, H., Bloomfield, G. S., & Sloth, E. (2020). Building Focused Cardiac Ultrasound Capacity in a Lower Middle-Income Country: A Single Centre Study to Assess Training Impact. *African Journal of Emergency Medicine*, 10(3), 136–143. <https://doi.org/10.1016/j.afjem.2020.04.011>
- Winter, M., Pryss, R., Probst, T., & Reichert, M. (2021). Applying Eye Movement Modeling Examples to Guide Novices' Attention in the Comprehension of Process Models. *Brain Sciences*, 11(1), 72. <https://doi.org/10.3390/brainsci11010072>
- Wolf, T., Sebanz, N., & Knoblich, G. (2018). Joint Action Coordination in Expert-Novice Pairs: Can Experts Predict Novices' Suboptimal Timing? *Cognition*, 178, 103–108. <https://doi.org/10.1016/j.cognition.2018.05.012>
- Zablocki, É., Piwowarski, B., Soulier, L., & Gallinari, P. (2018). Learning Multi-Modal Word Representation Grounded in Visual Context. *Proceedings of the Aaai Conference on Artificial Intelligence*, 32(1). <https://doi.org/10.1609/aaai.v32i1.11939>
- Piotrowska, B. M. (2024). Street-level Bureaucracy and Democratic Backsliding: Evidence From Poland. *Governance*, 37(S1), 127–151. <https://doi.org/10.1111/gove.12876>
- Schlechtweg, D., Hättig, A., Tredici, M. D., & Walde, S. S. i. (2019). A Wind of Change: Detecting and Evaluating Lexical Semantic Change Across Times and Domains. 732–746. <https://doi.org/10.18653/v1/p19-1072>
- Sevcik, N. (2022). Analyzing Democratic Backsliding in East Asia and the Implications for Human Rights. <https://doi.org/10.31235/osf.io/p7cjk>
- Shoemark, P., Liza, F. F., Nguyen, D., Hale, S. A., & McGillivray, B. (2019). Room to Glo: A Systematic Comparison of Semantic Change Detection Approaches With Word Embeddings. <https://doi.org/10.18653/v1/d19-1007>
- Weiss, M. (2022). Is Malaysian Democracy Backsliding or Merely Staying Put? *Asian Journal of Comparative Politics*, 9(1), 9–24. <https://doi.org/10.1177/20578911221136066>
- Wolkenstein, F. (2021). European Political Parties' Complicity in Democratic Backsliding. *Global Constitutionalism*, 11(1), 55–82. <https://doi.org/10.1017/s2045381720000386>

- Al-Thubaiti, K. (2018). Selective vulnerability in very advanced L2 grammars: evidence from vpe constraints. *Second Language Research*, 35(2), 225-252. <https://doi.org/10.1177/0267658317751577>
- Asudeh, A. (2022). Glue semantics. *Annual Review of Linguistics*, 8(1), 321-341. <https://doi.org/10.1146/annurev-linguistics-032521-053835>
- Bjorkman, B. (2022). Some structural disanalogies between pronouns and tenses. *The Canadian Journal of Linguistics / La Revue Canadienne De Linguistique*, 67(3), 143-165. <https://doi.org/10.1017/cnj.2022.26>
- Boleda, G. (2020). Distributional semantics and linguistic theory. *Annual Review of Linguistics*, 6(1), 213-234. <https://doi.org/10.1146/annurev-linguistics-011619-030303>
- Borgonovo, C., Garavito, J., & Prévost, P. (2014). Mood selection in relative clauses. *Studies in Second Language Acquisition*, 37(1), 33-69. <https://doi.org/10.1017/s0272263114000321>
- Bouveret, M. (2021). Lexicalization, grammaticalization and constructionalization of the verb give across languages., 1-22. <https://doi.org/10.1075/cal.29.int>
- Devine, A., & Stephens, L. (2017). Towards a syntax-semantics interface for Latin. *Catalan Journal of Linguistics*, 16, 79. <https://doi.org/10.5565/rev/catjl.210>
- Dudschig, C., Kaup, B., Liu, M., & Schwab, J. (2021). The processing of negation and polarity: an overview. *Journal of Psycholinguistic Research*, 50(6), 1199-1213. <https://doi.org/10.1007/s10936-021-09817-9>
- Ellis, N., O'Donnell, M., & Römer, U. (2014). Second language verb-argument constructions are sensitive to form, function, frequency, contingency, and prototypicality. *Linguistic Approaches to Bilingualism*, 4(4), 405-431. <https://doi.org/10.1075/lab.4.4.01ell>