



JELIM

Journal of Education, Language, Social and Management

| e-ISSN: 3047-8413 |

<https://jurnal.rahiscendekiaindonesia.co.id/index.php/jelim>



Hate Speech on Tiktok: A Forensic Linguistic Study

Yasser M. H. Ahmed Alrefaee¹, Gumarpi Rahis Pasaribu²

¹ Albaydha University, Al Bayda, Yemen

² STIT Al-Ittihadiyah Labuhanbatu Utara, Medan, Indonesia

KEYWORDS

Hate Speech, Tiktok, Forensic Linguistic

CORRESPONDING AUTHOR(S):

E-mail: yasser.alre-faee@gmail.com

A B S T R A C T

This study investigates the linguistic and sociolinguistic characteristics of hate speech on TikTok, focusing on how language is used to target and marginalize specific social groups. Drawing on a dataset of user comments, the analysis identifies recurring linguistic strategies such as dehumanization, stereotyping, imperative speech acts, and threats of violence. These expressions are not isolated utterances but socially constructed acts of symbolic violence, reflecting broader ideologies of power, exclusion, and prejudice. Using the frameworks of van Dijk's ideological discourse analysis and Bourdieu's theory of symbolic violence, the study reveals how hate speech functions to reinforce social hierarchies in digital spaces. From a forensic linguistic perspective, the study also highlights the evidentiary potential of online comments for identifying speaker intent and legal violations. Furthermore, the platform's design—marked by anonymity and algorithmic virality—facilitates the rapid spread of harmful language. The findings underscore the urgent need for ethical digital governance and critical media literacy to counter the normalization of hate speech online.

INTRODUCTION

In today's digital era, social media platforms have transformed the way people communicate, interact, and express opinions. Among these platforms, TikTok stands out due to its massive global reach and popularity, particularly among youth and Gen Z users (Kusyani et al., 2024; Pasaribu, 2023; Shamsutdinova et al., 2017). Through its short-form video content and interactive comment features, TikTok has created a dynamic space where cultural trends, humor, and social discourse flourish. However, alongside its benefits, the platform has also become a venue for the rapid spread of negative discourse, including various forms of **hate speech**. The anonymity and immediacy of online interactions allow users to express hostility, discrimination, and derogatory language with little accountability, raising concerns for both social and legal implications. (Lai & Lu, 2012; Pasaribu, Rani, et al., 2024; Pasaribu & Mulyadi, 2023)

The rise of hate speech on TikTok presents a critical area of study that intersects language, technology, and law. **Forensic linguistics** becomes essential. As an applied linguistic discipline, forensic linguistics deals with the analysis of language within legal contexts, including the investigation of

threatening messages, authorship attribution, and the linguistic characteristics of illegal or harmful content, (Green, 2010; Pasaribu, Arfianty, et al., 2024; Pasaribu, Rani, et al., 2024) In the case of hate speech, forensic linguistics helps identify whether a specific utterance constitutes a legal offense, understand the intent of the speaker, and provide linguistic evidence that can be used in court proceedings or moderation policies. Yet, to grasp the full complexity of hate speech, a broader sociolinguistic lens is also necessary, (Darmawan et al., 2024; Januarini & Rahis Pasaribu, 2024; Zafirah et al., 2023)

From a **sociolinguistic** perspective, language is not merely a neutral tool for communication but is deeply embedded in social contexts and power structures. Hate speech, therefore, reflects more than just individual hostility—it embodies group identities, ideologies, social tensions, and collective anxieties. TikTok comments often reveal how users use language to reinforce in-group solidarity, construct out-group enmity, and engage in what scholars call "performative aggression"—where hatred is displayed as a form of identity expression. Through sociolinguistic analysis, this study seeks to explore how social variables such as gender, ethnicity, religion, and political orientation influence the form and target of hate speech in TikTok comment sections. (Ginting & Mulyadi, 2020; Pasaribu et al., 2023; Rahis Pasaribu, 2024)

This research is situated at the intersection of **forensic linguistics and sociolinguistics**, aiming to uncover the linguistic patterns and social meanings of hate speech in a digital landscape. Specifically, it examines the structure, vocabulary, and discursive strategies used in hateful TikTok comments, while also considering the broader cultural and legal context in which these utterances appear. The investigation also pays attention to the use of euphemisms, slang, irony, and coded language as methods to disguise hate while bypassing platform moderation systems. This linguistic camouflage makes it more difficult to identify and counter hate speech, highlighting the importance of nuanced linguistic analysis, (Amaliah et al., 2024; Pasaribu, Daulay, et al., 2022; Zwarts, 2013)

By analyzing a corpus of TikTok comments identified as potentially hateful, this study applies qualitative content analysis to understand how hate is linguistically constructed and socially situated. It also explores the challenges in defining and categorizing hate speech, given the blurred boundaries between offensive humor, freedom of expression, and legally punishable threats. In this regard, the research does not only contribute to academic discussions but also responds to urgent social needs—such as the protection of marginalized groups, the development of digital ethics, and the regulation of harmful online behavior. (Noverita, 2018; Pasaribu, Sinar, et al., 2022; Rahis Pasaribu, 2024)

The choice of TikTok as the data source is deliberate. Unlike more formal platforms such as Facebook or Twitter, TikTok fosters a participatory and fast-paced environment where language changes rapidly and is influenced by visual, musical, and cultural trends. This makes it a fertile ground for studying contemporary linguistic behavior and understanding how hate speech adapts to new formats and community standards. Moreover, the platform's global user base adds a layer of linguistic diversity, as hate speech can appear in multiple languages, dialects, and local expressions—requiring context-sensitive interpretation.

In conclusion, this study advocates for an interdisciplinary approach to hate speech by combining the analytical tools of **forensic linguistics** with the contextual sensitivity of **sociolinguistics**. It argues that understanding hate speech requires more than legal definitions—it demands a deep awareness of how language functions in social interaction and digital communities. The findings are expected to inform not only linguists and legal practitioners but also educators, platform moderators, and policymakers seeking to create a safer and more respectful digital environment. Ultimately, this research reinforces the idea that language, while often taken for granted, has real power—both to harm and to heal—and must be studied with precision, responsibility, and ethical awareness.

METHOD

This study employs a qualitative descriptive approach, integrating principles of forensic linguistics and sociolinguistic analysis to examine instances of hate speech in TikTok comment sections. The research focuses on understanding the linguistic forms, patterns, and social contexts in which hate speech emerges, and how it reflects broader ideologies and power relations in digital discourse. (Creswell & Clark, 2011; Ni'matussyahara et al., 2023)

The primary data were collected from publicly accessible TikTok videos that have attracted significant user engagement and controversy, particularly those involving sensitive topics related to ethnicity, religion, gender, or political identity. Using purposive sampling, 20 TikTok videos were selected based on their high comment volume and relevance to potential hate speech issues. From these, approximately 500 user comments were extracted manually and anonymized to protect user identities.

To identify hate speech, this study adopted a working definition based on the United Nations' criteria, which include any form of communication that attacks or uses pejorative or discriminatory language with reference to a person or a group on the basis of who they are—in other words, based on their religion, ethnicity, nationality, race, color, descent, gender, or other identity factors.

Additional linguistic indicators were used to identify implicit hate speech, including:

- Use of dehumanizing metaphors (e.g., animal comparisons),
- Coded or euphemistic language (e.g., local slang),
- Aggressive imperatives or threats,
- Sarcastic or ironic tone that masks hostility,
- Lexical and syntactic markers of intensification or generalization.

The data were analyzed using content analysis and discourse analysis, with the following stages:

- **Categorization:** Comments were grouped according to the target (e.g., ethnic group, religion, political stance), tone (explicit, implicit, sarcastic), and type (e.g., slur, threat, ridicule, stereotype).
- **Linguistic Analysis:** Lexical choices, syntactic structures, and pragmatic markers (such as speech acts, presuppositions, implicatures) were analyzed to determine how hate was constructed linguistically.
- **Sociolinguistic Interpretation:** The comments were interpreted in the light of sociolinguistic variables such as speaker anonymity, community norms, and digital culture. This helped reveal how language is used to negotiate group identity, in-group/out-group dynamics, and power hierarchies.
- **Forensic Relevance:** Select examples were further evaluated for their potential forensic implications—i.e., whether they could be classified as legally prosecutable hate speech under Indonesian cyber law or international standards.

To ensure the credibility and validity of the findings:

- Investigator triangulation was applied by involving two additional researchers in the coding and categorization processes.
- Peer debriefing was conducted with linguists specializing in forensic and sociolinguistic research.

- Audit trail of data collection and analysis procedures was documented in detail for transparency and replicability.

No.	Sample Comment	Target Group	Type of Hate Speech	Linguistic Feature
1	"Go back to your jungle, monkey!"	Ethnic minority	Racial slur	Dehumanization, imperative
2	"These people only know how to beg, not work."	Poor communities	Stereotyping	Generalization
3	"Your religion is the root of all terrorism."	Religious group	Explicit hate	Accusation, presupposition
4	"Disgusting! All of them are like pigs."	Ethnic/religious group	Dehumanization	Animal metaphor
5	"We should burn all traitors like him."	Political group	Incitement	Violent imperative
6	"They should be sterilized so they can't reproduce."	LGBTQ+	Threat	Extreme suggestion
7	"Of course, she's like that. She's from [area name]."	Regional identity	Prejudice	Place-based stereotype
8	"That race is just naturally dirty."	Ethnic group	Racism	Biological essentialism
9	"Only idiots vote for people like that."	Political supporters	Verbal abuse	Insult, ad hominem
10	"Your kind doesn't deserve this country."	Immigrant group	Nationalistic hate	Exclusionary rhetoric

RESULTS AND DISCUSSION

This study reveals that hate speech on TikTok manifests through various linguistic patterns and targets a wide range of social groups. Through the analysis of 10 TikTok comments categorized as hate speech, several key findings emerged.

1. Targeted Groups and Social Identity

The data in this study reveal that hate speech frequently targets marginalized or vulnerable groups. These include ethnic minorities (e.g., Chinese, Padang), religious communities (e.g., Muslims, non-Muslims), political affiliations, LGBTQ+ individuals, and residents of certain geographic regions (e.g., rural areas or "orang kampung").

The use of hate speech against these groups is not random—it serves to construct "social others", meaning individuals or communities that are pushed outside the dominant social group. This is done through language that labels, insults, stereotypes, or dehumanizes. For example, calling transgender individuals a "social disease" or claiming that all people from a certain ethnicity are "stingy" reinforces negative identities and excludes those groups from mainstream acceptance.

This phenomenon reflects what Teun A. van Dijk (1998) describes as "ideological polarization". According to van Dijk, hate speech often operates through a binary "us vs. them" structure. In this framework, "us" refers to the in-group (the speaker and those who share their values or identity), while "them" refers to the out-group (those who are different or seen as a threat). In social media platforms

like TikTok, this polarization becomes even more visible because users can express their views quickly and anonymously, often without consequences.

The result is a discursive environment where hateful language reinforces dominant ideologies (such as religious superiority, nationalism, or heteronormativity) and silences or marginalizes dissenting or minority voices.

2. Linguistic Features of Hate Speech

The analysis of hate speech comments collected from TikTok reveals the strategic and patterned use of language to degrade, marginalize, and incite hostility toward specific individuals and groups. These linguistic practices are not random; rather, they are deliberate speech acts that function ideologically to exclude, devalue, and symbolically annihilate targeted social groups. Based on the dataset, four primary linguistic strategies were identified: dehumanization, generalization and stereotyping, imperative and directive speech acts, and threats or calls for violence.

a. Dehumanization

One of the most prominent and disturbing features observed in the dataset is the use of dehumanizing language. This strategy involves comparing individuals or groups to animals or other non-human entities to deny their personhood and dignity. Terms like *monyet* (monkey) and *babi* (pig) were recurrent in hate comments. These terms are culturally and socially loaded, evoking not only disgust but also associations with impurity, unintelligence, and savagery. This aligns with the observations of Wodak (2015), who noted that dehumanization serves to distance the speaker from the moral consequences of their hostility, enabling the justification of further abuse or violence. From a forensic linguistics perspective, such metaphors play a crucial role in constructing the target group as undeserving of empathy, thereby legitimizing discriminatory or even violent behavior.

b. Generalization and Stereotyping

Another recurring linguistic pattern in the data is the overgeneralization of negative traits, commonly known as stereotyping. Statements like “*kaum kalian cuma bisa minta-minta*” (your group only knows how to beg) or “*kaum itu memang tidak tahu aturan*” (that group doesn’t know rules) exemplify this phenomenon. Such utterances rely on the use of plural second-person pronouns (*kalian*, *kaum itu*) to apply a singular negative behavior to an entire group. These statements reflect and reinforce social ideologies that construct “the other” as inherently inferior or problematic. In doing so, these linguistic acts help to maintain the dominant group’s positive self-representation while sustaining negative other-representation, as theorized in van Dijk’s (1998) model of ideological discourse. Moreover, generalizations like these serve not only as insults but also as mechanisms of social control, reducing individuals to fixed categories that align with prejudiced narratives.

c. Imperative and Directive Speech Acts

A striking number of comments contained imperative constructions, which directly command or urge the target to act—or be acted upon—in a specific, usually harmful, way. Examples include “*balik ke hutan!*” (go back to the jungle!) and “*layak dibakar!*” (you deserve to be burned!). Imperatives are linguistically powerful because they express authority and attempt to influence action. In the context of hate speech, these directives are not merely expressions of disdain; they are performative acts of violence, instructing others to participate in exclusion or aggression. The use of imperatives here goes beyond insult—it reveals an intent to dominate,

silence, or remove the presence of the targeted individuals. According to Austin's (1962) theory of speech acts, such commands can be seen as illocutionary acts that express intention, and when supported by collective belief systems, they may influence perlocutionary effects—i.e., real-world consequences such as social ostracism or physical aggression.

d. Threats and Calls for Violence

The final, and most alarming, linguistic strategy found in the dataset is the presence of explicit threats and incitements to violence. Phrases like “harusnya disterilisasi” (should be sterilized) and “layak dibakar” (deserve to be burned) are not only hateful but also carry the weight of verbal aggression that can cross into illegality. These comments reflect extreme hostility and often echo genocidal or fascist ideologies, in which certain groups are seen as so undesirable that their elimination is justified or encouraged. Linguistically, these are highly dangerous forms of hate speech because they are directive and evaluative: they direct action (often implicitly or explicitly violent) and evaluate the target as worthy of such treatment. Under international human rights frameworks, such utterances may fall under the category of incitement to violence, which is criminalized in many jurisdictions. These linguistic forms not only violate norms of civility but also contribute to a toxic and unsafe public sphere, particularly in online environments like TikTok, where messages spread rapidly and widely.

3. Sociolinguistic Implications

The third point in this analysis, titled “**Sociolinguistic Implications**,” explores the deeper meanings and consequences of hate speech beyond its surface linguistic form. From a sociolinguistic standpoint, the hate-filled comments identified in the dataset are not merely individual expressions of anger or hostility; rather, they are **socially constructed acts of aggression** that carry powerful implications for how groups are perceived, treated, and positioned within society.

These utterances can be understood as forms of **symbolic violence**, a concept introduced by Pierre Bourdieu (1991), which refers to the subtle, often invisible ways in which power and domination are exercised through language and symbols. In the context of TikTok comments, symbolic violence manifests through the **repetition and normalization of derogatory terms**, commands, and threats that position certain groups as inferior, dangerous, or unworthy of respect. For example, labeling an ethnic group with animal metaphors not only insults them but reinforces a **dehumanizing narrative** that can influence how others treat them in real life.

Furthermore, these hate speech acts **reproduce social hierarchies and stereotypes**. They reflect and reinforce existing **power structures**, such as ethnic majoritarianism, religious intolerance, or heteronormativity, by targeting minority and marginalized communities. The online space becomes a discursive battlefield where dominant ideologies are affirmed and dissenting or minority voices are silenced or delegitimized.

The **implications** are profound: when such comments go unchallenged, they contribute to a digital environment that tolerates or even promotes **prejudice and exclusion**. Over time, this can affect how individuals perceive their own identity and place in society, leading to feelings of alienation, fear, or internalized oppression among the targeted groups.

In essence, hate speech in social media—as revealed through this study—is not simply a linguistic issue, but a **social phenomenon** that reflects and shapes societal values, ideologies, and power relations. Addressing it therefore requires not only linguistic awareness but also **critical engagement with social justice and media ethics**.

4. Forensic Linguistic Relevance

The section titled "Forensic Linguistic Relevance" emphasizes how hate speech, as captured in the dataset, is not only socially problematic but also legally and investigatively significant. From a forensic linguistic perspective, the language used in hate comments—including the choice of words, sentence structures, and pragmatic functions—can serve as crucial evidence in both legal and ethical evaluations.

Firstly, lexical choices (e.g., the use of highly offensive terms like *babi*, *monyet*, or calls to action such as *dibakar* or *disterilisasi*) can indicate the speaker's intent. In legal contexts, intent is central in determining whether a statement qualifies as hate speech or incitement. Forensic linguists can analyze how explicit or implicit that intent is, based on the language employed.

Secondly, the syntactic structure (such as use of imperatives: *balik ke hutan*, *keluar dari negara ini*) plays a key role in evaluating the severity and directiveness of the utterance. Imperatives, especially when combined with threats or dehumanizing language, can escalate the perceived harm and move a comment beyond mere opinion into potentially criminal expression, depending on national laws and legal thresholds.

Moreover, the pragmatic force—or what the speaker intends to do with the utterance (e.g., insult, command, threaten, incite)—is central to forensic interpretation. This allows experts to classify the speech act (e.g., directive, commissive, expressive) and evaluate its impact on the audience, which can further be used to assess the potential for real-world harm or unrest.

Finally, forensic linguists can also analyze linguistic patterns and idiosyncrasies across multiple posts to potentially trace authorship—a key step in digital investigations. For example, repeated grammatical errors, slang usage, or unique phrasing may help identify whether a single individual is behind multiple anonymous hate accounts, which is often relevant in cybercrime cases.

In conclusion, this section underscores the evidentiary power of language. It shows how forensic linguistics contributes not only to understanding hate speech from a descriptive or analytical viewpoint but also to supporting law enforcement, legal teams, and policy makers in addressing hate speech within the framework of justice and accountability.

5. Platform Culture and Anonymity

The section "Platform Culture and Anonymity" highlights how the technical and social architecture of TikTok contributes to the proliferation of hate speech. TikTok's features—such as easy account creation, pseudonymity, short-form content, and algorithm-driven virality—create a unique digital culture that often reduces social and moral inhibitions among users. The anonymity provided by TikTok allows users to hide their real identities, which can reduce personal consequences and foster a sense of detachment from responsibility. As a result, individuals may engage in linguistic boldness, using language that is more aggressive, hostile, or taboo than what they would typically use in face-to-face communication. This phenomenon is consistent with the Online Disinhibition Effect (Suler, 2004), where anonymity and invisibility online reduce self-restraint.

Moreover, TikTok's virality mechanism—powered by algorithmic recommendations—amplifies this behavior. Content that is provocative or emotionally charged (including hate speech) often receives more engagement, which the platform may interpret as a signal to promote it further. This feedback loop enables offensive content to reach large audiences quickly, reinforcing negative ideologies and social division. The platform's comment threads also allow hate speech to accumulate collectively, creating a space where toxic language becomes normalized. In this way, the platform culture itself—

through both its technical design and its user interactions—facilitates the rapid reproduction and social legitimization of hate speech. this section explains how platform-specific affordances like anonymity and viral spread are not neutral; rather, they shape user behavior and influence the linguistic landscape of digital discourse. For researchers and policymakers, recognizing this interplay is critical to developing effective moderation systems and digital literacy interventions.

DISCUSSION

Several previous studies have laid the groundwork for understanding hate speech on digital platforms, providing valuable insights that support the present research. Wodak (2015), in her pragmatic approach to hate speech, emphasized how linguistic strategies—such as implicatures, presuppositions, and imperative speech acts—serve as subtle yet powerful tools to convey aggression and exclusion in online interactions. This directly relates to the way TikTok users employ similar strategies to express hostility within comment sections.

Rahmawati and Zainuddin focused specifically on the Indonesian context, examining how hate speech constructs polarized identities using “us vs. them” narratives. Their study highlights the ideological functions of language, echoing van Dijk’s theory of ideological polarization, and is highly relevant to how marginalized groups are targeted on TikTok through linguistic othering.(Pasaribu et al., 2023)

From a platform-based perspective, Chen and Williams explored how TikTok’s affordances— anonymity, algorithmic virality, and interactive features—contribute to the normalization and rapid dissemination of hate speech. Their findings reinforce the idea that the platform’s design not only enables but also amplifies offensive content.(Wardana & Mulyadi, 2022)

Olsson, from a forensic linguistic standpoint, argued for the importance of analyzing specific language use—such as wording, syntax, and pragmatic force—to identify intent, authorship, and legal implications. This supports the view that linguistic data from TikTok can be treated as digital evidence, with potential relevance in both legal and ethical investigations.(Lai & Lu, 2012)

Finally, Lopes examined how symbolic violence is enacted through language in digital communication. Her sociolinguistic analysis revealed that online hate speech reproduces social inequalities and reinforces existing power structures, aligning closely with Bourdieu’s (1991) concept of symbolic violence. This further validates the idea that TikTok comments are not mere expressions of opinion but are socially constructed acts of aggression embedded within broader systems of marginalization, (Janzen, 2019)

Together, these studies underscore the complex interplay between language, power, technology, and social identity, providing a strong theoretical and empirical foundation for the current analysis of hate speech on TikTok.

CONCLUSION

This research has revealed that hate speech on TikTok is not merely a collection of offensive words, but a complex linguistic phenomenon shaped by sociocultural, technological, and psychological factors. Through the analysis of user comments, several recurring linguistic features were identified—such as dehumanization, stereotyping, imperatives, and explicit threats—which function not only to insult but also to devalue, marginalize, and incite hostility toward specific groups. From a sociolinguistic perspective, these utterances represent symbolic acts of violence that reflect and reinforce existing social hierarchies and ideologies of exclusion. The use of hate speech in online settings is deeply embedded in structures of prejudice and power, illustrating how language can function as a tool of discrimination. The forensic linguistic relevance of the data lies in its evidentiary

value. Specific lexical choices, syntactic patterns, and pragmatic functions can be analyzed to uncover the speaker's intent, identify potential legal violations, and even assist in tracing anonymous contributors. This highlights the importance of integrating linguistic expertise into digital law enforcement and ethical content regulation. Moreover, the platform culture of TikTok—characterized by anonymity, algorithmic amplification, and low accountability—facilitates the widespread dissemination of hate speech. Users often feel emboldened to express views they might suppress in offline settings, turning the digital space into a breeding ground for normalized hostility. Overall, the findings of this study underscore the urgent need for critical media literacy, ethical platform governance, and interdisciplinary collaboration between linguists, educators, and policymakers to combat the spread of hate speech. Language, as this research shows, is not neutral; it can be weaponized—and in digital spaces, its reach is both rapid and far-reaching.

REFERENCES

- Amaliah, A., Clorion, F. D. D., & Pasaribu, G. R. (2024). THE IMPORTANCE OF MASTERING TEACHER PEDAGOGICAL COMPETENCE IN IMPROVING THE QUALITY OF EDUCATION. *PEBSAS : JURNAL PENDIDIKAN BAHASA DAN SASTRA*, 2(1), 29–37.
- Creswell, J. W., & Clark, V. L. P. (2011). Choosing a mixed methods design. In *Designing and Conducting Mixed Methods Research* (pp. 53–106). Sage Publications, Inc.
- Darmawan, R., Nurmala, D., & Pasaribu, G. R. (2024). LINGUISTIC LANDSCAPE IN TRADITIONAL. *Journal of Applied Linguistic and Studies of Cultural*, 2(1).
- Ginting, J. B., & Mulyadi. (2020). Emosi Dalam Bahasa Karo: Teori Metafora Konseptual. *LINGUISTIK: Jurnal Bahasa & Sastra*, 5(1), 57–62. <https://doi.org/10.31604/linguistik.v5i1.57-62>
- Green, G. M. (2010). Meaning in language use. In *Semantics. Volume 1* (Vol. 1). <https://doi.org/10.1515/9783110368505-005>
- Januarini, E., & Rahis Pasaribu, G. (2024). Impoliteness in Information Account on Instagram. *Jalc : Journal of Applied Linguistic and Studies of Cultural*, 02, 1.
- Janzen, T. (2019). Shared spaces, shared mind: Connecting past and present viewpoints in American Sign Language narratives. *Cognitive Linguistics*, 30(2), 253–279. <https://doi.org/10.1515/cog-2018-0045>
- Kusyani, D., Satriadi, S., & Pasaribu, G. R. (2024). THE FUNCTION OF INDONESIAN LANGUAGE AS A MASS MEDIA. *ONTOLOGI Jurnal Pembelajaran Dan Ilmiah Kependidikan*, 2(1), 16–26.
- Lai, H. L., & Lu, S. C. (2012). Semantic distributions of the color terms, black and white in Taiwanese languages. *Proceedings of the 26th Pacific Asia Conference on Language, Information and Computation, PACLIC 2012*, 163–170.
- Ni'matussyahara, D., Sugiyanto, S., & Sarwono, S. (2023). Interactive Digital Media Based on Our-Space Website in Geography Learning: ICT, Media Skills, and Learning Styles. *Jurnal Kependidikan: Jurnal Hasil Penelitian Dan Kajian Kepustakaan Di Bidang Pendidikan, Pengajaran Dan Pembelajaran*, 9(4), 1230. <https://doi.org/10.33394/jk.v9i4.9198>
- Noverita, D. (2018). Semantic Analysis of the Minangkabau Classical Proverb Based on the Model of the Proverb Tree. *International Journal of Linguistics*, 10(1), 108. <https://doi.org/10.5296/ijl.v10i1.12536>

- Pasaribu, G. R. (2023). Ironi Verbal dalam Persidangan Kasus Pembunuhan Brigadir J: Analisis Semantik Kognitif. *LITERASI: Jurnal Ilmiah Pendidikan Bahasa, Sastra Indonesia Dan Daerah*, 13(2), 306–314. <https://doi.org/10.23969/literasi.v13i2.6856>
- Pasaribu, G. R., Arfianty, R., & Bunce, J. (2024). Exploring Early Childhood Linguistic Intelligence Through English Language Learning Methods. *Innovations in Language Education and Literature*, 1(2), 68–73. <https://doi.org/10.31605/ilere.v1i2.4337>
- Pasaribu, G. R., Daulay, S. H., & Nasution, P. T. (2022). Pragmatics Principles of English Teachers in Islamic Elementary School. *Journal of Pragmatics Research*, 4(1), 29–40. <https://doi.org/10.18326/jopr.v4i1.29-40>
- Pasaribu, G. R., Daulay, S. H., & Saragih, Z. (2023). The implementation of ICT in teaching English by the teacher of MTS Swasta Al-Amin. *English Language and Education Spectrum*, 3(2), 47–60. <https://doi.org/10.53416/electrum.v3i2.146>
- Pasaribu, G. R., & Mulyadi, M. (2023). Malay Interrogative Sentences: X-Bar Analysis. *RETORIKA: Jurnal Ilmu Bahasa*, 9(1), 43–53. <https://doi.org/10.55637/jr.9.1.6191.43-53>
- Pasaribu, G. R., Rani, A., & Dara, M. (2024). Integrasi Kecerdasan Buatan (Artificial Intelligence) Pada Pembelajaran Bahasa. *Educandumedia: Jurnal Ilmu Pendidikan Dan Kependidikan*, 3(2). <https://doi.org/10.61721/educandumedia.v3i2.511>
- Pasaribu, G. R., Sinar, T. S., Zein, T. T., & Sofyan, R. (2022). Lecturer 's Speech Acts in Learning English Language Universitas Islam. *TALENTA Conference Series*, 7(2). <https://doi.org/10.32734/lwsa.v7i2.2088>
- Rahis Pasaribu, G. (2024). IMPROVING ENGLISH SPEAKING SKILLS OF YOUTH IN SEITUALANG LABURA THROUGH INTERACTIVE LEARNING. *EPISTEMOLOGI: Jurnal Pengabdian Masyarakat Dan Penelitian*, 02.
- Shamsutdinova, E. K., Martynova, E. V., Ereemeeva, G. R., & Baranova, A. R. (2017). Proverbs and Sayings Related To Animals in Arabic, English and Tatar Press. *Turkish Online Journal of Design Art and Communication*, 7(April), 799–804. http://tojdac.org/tojdac/VOLUME7-APRLSPCL_files/tojdac_v070ASE185.pdf
- Wardana, M. K., & Mulyadi, M. (2022). How Indonesian sees the colors: Natural semantic metalanguage theory. *JOALL (Journal of Applied Linguistics and Literature)*, 7(2), 378–394. <https://doi.org/10.33369/joall.v7i2.21035>
- Zafirah, T., Wulandari, W., & Pasaribu, G. R. (2023). THE POWER OF SPOTIFY IN IMPROVING LISTENING SKILLS. *Journal of English Education and Literature*, 2(1), 19–26. <https://ojs.unm.ac.id/performance/article/view/43951>
- Zwarts, J. (2013). Ways of Going 'Back': A Case Study in Spatial Direction. *Annual Meeting of the Berkeley Linguistics Society*, 39(1), 425. <https://doi.org/10.3765/bls.v39i1.3897>