



Hate Speech in the Comments Section of the *Hilirisasi* Video Series on the @GibranTV YouTube Channel: A Forensic Linguistic Study

Aida^{1*}, Riki Nasrullah²

¹²Indonesian Language and Literature Education Study Program, Surabaya State University, East Java, Indonesia

*E-mail: aida.2020074083@mh.s.unesa.ac.id

ABSTRACT

This study aims to describe hate speech in the comments section of the YouTube series "Hilirisasi" on the @GibranTV channel through a forensic linguistic study using a pragmatics scalpel. The focus of the study includes the forms of hate speech, types of illocutionary speech acts, violations of politeness maxims, and forensic linguistic categories that appear in netizens' comments. The method used is descriptive qualitative, with data in the form of netizens' comments containing indications of hate speech in the video's comment section. Hate speech data were collected using observation guidelines based on time triangulation. The research data were analyzed thematically. The results show that hate speech in the comment section is manifested in the form of insults, accusations, negative labeling, invitations, and provocation. The speech is dominated by expressive, assertive, and directive speech acts, and contains violations of the maxims of praise, sympathy, wisdom, agreement, generosity, and simplicity. From a forensic linguistic perspective, the speech found can be categorized as insults, provocation, incitement, defamation, and indications of the dissemination of unverified information. These findings suggest that digital comments not only serve as a means of expressing opinions, but also have the potential to become a form of verbal aggression.

Keywords: hate speech, comments section, video series, Youtube channel, forensic linguistics

Ujaran Kebencian pada Kolom Komentar Seri Video Bertema *Hilirisasi* di Akun Youtube @GibranTV: Kajian Linguistik Forensik

ABSTRAK

Penelitian ini bertujuan untuk mendeskripsikan ujaran kebencian dalam kolom komentar YouTube seri "Hilirisasi" pada kanal @GibranTV melalui kajian linguistik forensik menggunakan pisau bedah kajian pragmatik. Fokus penelitian mencakup bentuk ujaran kebencian, jenis tindak tutur ilokusi, pelanggaran maksim kesantunan, serta kategori linguistik forensik yang muncul dalam komentar warganet. Metode yang digunakan berupa deskriptif kualitatif, dengan data berupa komentar warganet yang mengandung indikasi ujaran kebencian pada kolom komentar video tersebut. Data ujaran kebencian dikumpulkan menggunakan pedoman observasi berbasis triangulasi waktu. Data penelitian dianalisis secara tematik. Hasil penelitian menunjukkan bahwa ujaran kebencian dalam kolom komentar diwujudkan dalam bentuk penghinaan, tuduhan, pelabelan negatif, ajakan, dan provokasi. Ujaran tersebut didominasi oleh tindak tutur ekspresif, asertif, dan direktif, serta mengandung pelanggaran maksim pujian, simpati, kearifan, kesepakatan, kedermawanan, dan kesederhanaan. Dalam perspektif linguistik forensik, tuturan yang ditemukan dapat dikategorikan sebagai penghinaan, provokasi, hasutan, pencemaran nama baik, dan indikasi penyebaran informasi yang belum terverifikasi. Temuan ini menunjukkan bahwa komentar digital tidak hanya berfungsi sebagai sarana berpendapat, tetapi juga berpotensi menjadi bentuk agresi verbal.

Kata kunci: ujaran kebencian, kolom komentar, akun Youtube, linguistik forensik

Submitted
07/07/2026

Accepted
22/07/2026

Published
27/07/2026

Citation	Aida, A., & Nasrullah, N. (2026). Hate Speech in the Comments Section of the <i>Hilirisasi</i> Video Series on the @GibranTV YouTube Channel: A Forensic Linguistic Study. <i>Jurnal Pembelajaran Bahasa dan Sastra</i> , Volume 5, Nomor 4, Juli 2026, 3323-3334. DOI: https://doi.org/10.55909/jpbs.v5i4.1715
----------	---

Publisher
Raja Zulkarnain Education Foundation

INTRODUCTION

The development of information and communication technology has transformed the patterns of social interaction in conveying information, ideas, and opinions in the digital public sphere. One aspect that has undergone changes due to this development is the use of language in public communication (Syafuruddin, Thaba, A. & Ananda, 2024:6). The presence of social media allows every individual to actively participate in various discussions without being limited by space and time. This freedom makes social media function not only as a means of sharing information but also as a space for forming public opinion on various issues and politics.

One platform widely used by the public to express opinions is YouTube through its comment section feature, which allows for two-way communication between content creators and other users. YouTube is a popular social media platform today as a source of information and communication (Ernawati & Prayitno, 2024:71). This means that YouTube is a social media platform for information and communication today.

The comment section on YouTube has evolved into a discussion space that displays diverse public responses to information or events. Users can express appreciation, criticism, suggestions, or even rejection of the content of the videos shown. However, this freedom of expression is not always realized through the use of polite language.

Digital footprints are permanent and can serve as authentic evidence in legal proceedings (Mintowati, Nasrullah, & Ahmadi, 2025:16). This situation means that every comment published in the digital space has the potential to have legal consequences. Digital footprints not only record online interactions but can also serve as valid evidence in court. User anonymity, minimal social control, and the high intensity of debate in the digital space encourage various forms of aggressive communication, such as insults, negative labeling, accusations, provocation, and hate speech against certain individuals or groups.

In this context, hate speech is understood as linguistic acts aimed at attacking, demeaning, inciting, or inciting hatred against individuals or groups based on their identity, views, or status. In the digital realm, hate speech is not simply an outburst of emotion but also a communication tactic that reflects worldviews, resistance, and power dynamics in online public interactions (Mintowati, Nasrullah, & Ahmadi, 2025:45-47). Hate speech is not simply defined as the use of harsh words, but rather as a linguistic act aimed at attacking, demeaning, inciting, or inciting hatred against individuals or groups based on their identity, views, or status. Based on the Circular Letter of the Chief of Police Number SE/6/X/2015, hate speech can be manifested in the form of insults, defamation, blasphemy, provocation, incitement, the spread of false news, or other actions that have the potential to cause discrimination, violence, or social conflict. Thus, hate speech is a linguistic issue that has both social and legal consequences.

From a legal perspective, the dissemination of hate speech through electronic media receives special attention through Law Number 1 of 2024 concerning the Second Amendment to Law Number 11 of 2008 concerning Electronic Information and Transactions. This regulation serves as the legal basis for handling the dissemination of electronic information containing insults, defamation, or incitement that incites hatred or hostility. These acts can also be prosecuted under the provisions of the Criminal Code and the Electronic Information and Transactions Law if carried out through electronic media (Mintowati, Nasrullah, & Ahmadi, 2025:74). However, the application of these legal provisions remains controversial because not all forms of criticism or expressions of disapproval can be categorized as hate speech. Therefore, a linguistic analysis is needed that can explain the meaning, intent, and context of a speech's use, allowing for a distinction between legitimate criticism and potentially unlawful speech.

One approach that can be used to address this issue is forensic linguistics. In forensic linguistics,



language is viewed as evidence that can be scientifically analyzed to reveal the speaker's intent, the form of the speech, and its legal implications. According to McMenamin (2002:87), forensic linguistics is the study of language used in various legal contexts, including constitutional, statutory, and civil law. Every utterance leaves a linguistic trace that reflects the speaker's intent (*mens rea*). Recent research shows that hate speech, toxic language, and cyberbullying are forms of language that need to be addressed as soon as possible through a forensic linguistic framework (Mintowati, Nasrullah, & Ahmadi, 2025:23). The development of digital communication demands an expansion of forensic linguistic studies to explain the phenomena of hate speech, toxic language, and cyberbullying that are increasingly prevalent on various online platforms. Analysis should not only focus on the linguistic elements that appear on the surface but also consider the social context, relationships between speakers, and the purpose of the communication. Through this approach, hate speech can be examined more objectively, resulting in a basis for analysis that can be justified both linguistically and legally.

The analysis of hate speech also requires a pragmatic approach, specifically Leech's theory of illocutionary speech acts. The pragmatic approach views every utterance not only as conveying information but also as representing the speaker's actions through language. Pragmatics is the study of meaning in relation to speech situations (Leech, 2015:8). Therefore, the meaning of an utterance can only be understood by considering the relationship between the utterance and the context in which it is used. By identifying the type of illocutionary speech act used, the speaker's intended purpose can be determined when delivering a comment. Searle, in Leech (1969:164), classifies illocutionary speech acts into five types: assertive (describing), directive (giving orders), commissive (promising), expressive (expressing feelings), and declarative (stating a change in status). Therefore, speech act analysis is an important foundation for identifying the

characteristics of hate speech appearing on social media.

Hate speech appearing on social media also reflects violations of the principles of politeness. Therefore, this study also integrates the theory of violations of politeness maxims. The politeness maxims proposed by Leech (2015:206) provide guidelines for how speakers should act in interactions to maintain harmonious social relations. However, in communication practice, not all utterances adhere to these maxims. When speakers disregard the politeness maxims, they violate them.

This phenomenon is evident in the comments section of the video series "Downstreaming" uploaded to the YouTube account @GibranTV. The videos, which discuss the downstreaming policy, elicited a variety of responses from the public, ranging from support to criticism, conveyed in emotionally charged language. These comments included insults, negative labels, accusations, provocations, and other forms of speech that could potentially be categorized as hate speech. This situation indicates that the YouTube comment section is not only a space for expressing opinions but also an arena for the use of language that has the potential to have legal consequences. This phenomenon is interesting to study from a forensic linguistics perspective, combining analysis of illocutionary speech acts, the classification of hate speech based on Circular Letter of the Chief of Police Number SE/6/X/2015, and its legal implications under Law Number 1 of 2024 concerning Information and Electronic Transactions.

Based on this background, the research problem formulation in this study includes three questions. First, what types of illocutionary speech acts are used in the comments section of the video series "Downstreaming" on the YouTube account @GibranTV? Second, how are the types of violations of politeness maxims found in hate speech in the comments section of the "Hilirisasi" video series on the @GibranTV YouTube account classified? Third, how are the forms of hate speech

in the comments section of the "Hilirisasi" video series on the @GibranTV YouTube account classified based on the Chief of Police Circular Letter Number SE/6/X/2015 and their legal implications under Law No. 1 of 2024 concerning Electronic Information and Transactions (UU ITE)?

This study aims to describe the forms of illocutionary speech acts in the comments section of the "Hilirisasi" video series on the @GibranTV YouTube account, classify these forms of hate speech based on the Chief of Police Circular Letter Number SE/6/X/2015, and explain their legal implications under Law No. 1 of 2024 concerning Electronic Information and Transactions. This research is expected to provide theoretical and practical benefits. Theoretically, this research is expected to enrich forensic linguistic studies, particularly regarding illocutionary speech acts, hate speech, and their legal implications on social media. Practically, this research is expected to serve as a reference for researchers, academics, law enforcement officials, and the public in understanding and identifying hate speech in the digital space.

Extensive research has been conducted on hate speech on social media. Batubara and Mulyadi (2023) examined hate speech using a forensic linguistic approach and found a predominance of expressive illocutionary speech acts. Widyatnyana et al. (2023) analyzed forms of hate speech against political figures on social media, emphasizing linguistic aspects. Meanwhile, Manurung et al. (2021) identified hate speech based on linguistic elements containing negative emotions. Unlike previous studies, which tended to focus on linguistic aspects or speech acts separately, this study integrates forensic linguistic analysis, illocutionary speech act theory, the classification of hate speech based on the Circular Letter of the Chief of Police Number SE/6/X/2015, and an analysis of the legal implications under Law Number 1 of 2024 concerning Electronic Information and Transactions within a single analytical framework.

METHOD

This research employs a descriptive method with a qualitative approach. Razak (2017:176) dan Yuliani (2018:87) menyebutkan metode deskriptif lazim digunakan dalam penelitian bidang sosial seperti bahasa dan sastra. This method is used to describe the forms of illocutionary speech acts, violations of politeness maxims, and the classification of hate speech in the comments section of the "Hilirisasi" video series on the @GibranTV YouTube account, while also examining its legal implications under Law Number 1 of 2024 concerning Electronic Information and Transactions. A qualitative approach was chosen because it allows researchers to interpret the meaning of speech contextually based on digital communication situations.

The research data source consists of netizen comments on the "Hilirisasi" video series, parts 1 to 4, on the @GibranTV YouTube account, totaling approximately 31,713 comments. The research data focused on selected speech acts using purposive sampling, which selects data based on specific criteria relevant to the research objectives, namely comments containing illocutionary speech acts and indications of hate speech.

Data collection was conducted using the Free Listening and Conversation (SBLC) technique and note-taking. Data was obtained through the process of downloading comments using YouTube Comment Downloader, then manually verified to ensure completeness and validity. Next, the data was selected, classified, and coded according to the analysis requirements. The primary research instrument was the researcher, assisted by a data classification table and comment screenshot documentation.

The data analysis technique used qualitative content analysis. The analysis was conducted through the stages of coding, classification, and interpretation. During the analysis stage, data were identified based on three main aspects: types of illocutionary speech acts (Searle), violations of politeness maxims (Leech), and forms of hate speech based on Circular Letter of the Chief of



Police Number SE/6/X/2015. The analysis results were then linked to legal provisions in the ITE Law to examine the legal implications of the speech found.

Data validity was tested through source triangulation by comparing comments across the entire video series and expert judgment by a linguist to ensure the accuracy of the classification of speech acts, violations of politeness, and hate speech, ensuring scientific accountability of the research results.

RESULTS

1. Types of Illocutionary Speech Acts

Three types of illocutionary speech acts were found in hate speech in the comments section of the video series with the theme "downstreaming" on the YouTube account @GibranTV: assertive, directive, and expressive speech acts. Commissive and declarative speech acts were not found in the research data. Findings regarding each type of illocutionary speech act are presented based on the categories found in the research data.

1.1 Assertive Speech Acts

Assertive speech acts are found in speech acts that convey claims or assessments of the target of the hate speech. Speech in this category is positioned as statements believed to be true by the speaker. One piece of data representing assertive speech acts is shown below.

Data 1

"Likes are bought, let alone votes..." (V1/004). Data 1 is categorized as an assertive speech act because it contains a claim or assessment of the alleged actions of the target of the speech. This claim is conveyed in the form of a statement believed to be true by the speaker and meets the characteristics of an assertive illocutionary speech act. The use of the phrase "likes are just bought" indicates an accusation delivered as a statement. This utterance is then reinforced by the phrase "apalagi suara" (what else sounds), which develops the speaker's suspicions about the target of the

utterance.

1.2 Directive Speech Acts

The use of directive illocutionary acts in the research data is demonstrated through utterances containing commands, invitations, or encouragement to certain parties to perform an action. The presence of directive forms indicates that hate speech is not only expressed in the form of judgment, but also through attempts to direct the behavior of the target of the utterance. These characteristics are evident in the following data.

Data 2

"Stop pretending to be confident, above hypocrisy." (V2/019). The word "stop" is the primary marker indicating the directive function in the utterance. The command is directed at the target to stop behavior deemed to be pretense. The phrase "above hypocrisy" reinforces the reasoning behind the command. These linguistic characteristics indicate that this data is a directive speech act.

1.3 Expressive Speech Acts

Expressive speech acts were the most dominant form in the research data. This dominance indicates that speakers more frequently expressed negative attitudes and feelings toward the target through hate speech. Expressive forms are manifested in expressions such as insults, ridicule, sarcasm, and curses directed at the target. One form of expressive speech act is presented in the following data.

Data 3

"I feel nauseous seeing Gibran's face" (V1/010). This data demonstrates an expressive speech act because it implicitly expresses disgust directed at the target. The use of the word "nauseous" represents the speaker's emotional response to seeing the target. The utterance does not convey a command or claim, but rather directly expresses negative feelings. These characteristics indicate that this data is an expressive illocutionary speech act.

2. Types of Violations of Politeness Maxims

Violations of politeness maxims in the comments section of the video series "Downstreaming" on the YouTube account @GibranTV do not encompass all of the maxims proposed by Leech. The research data only revealed three types of violations: the maxim of appreciation, the maxim of agreement, and the maxim of sympathy. Each violation has distinct characteristics depending on the language used in hate speech.

2.1 Violations of the Maxim of Appreciation

Violations of the maxim of appreciation were the most common form found in the research data. This violation is characterized by the use of speech that demeans, insults, mocks, or negatively labels the target. These characteristics indicate that hate speech is more often manifested through attacks on individual dignity than through criticism of ideas or actions.

Data 4

"A burden on the state; its presence is not desired by the people. A trusted person, it's ridiculous." (V1/007)

Data 4 shows a violation of the maxim of agreement because it contains negative labeling through the expressions "a burden on the state, a trusted person, and very ridiculous." This utterance does not honor the target, but instead constructs a derogatory image through word choice with negative evaluative value. This data is categorized as a violation of the maxim of agreement.

2.2 Violation of the Maxim of Agreement

Violations of the maxim of agreement occur in utterances that demonstrate rejection of another party's opinion or attitude. This form of violation is characterized by the use of utterances that emphasize opposition without opening up space for reaching an understanding. This pattern demonstrates the speaker's tendency to maintain a

position that contradicts the target's. One type of violation of the maxim of agreement is shown in the following data.

Data 5

"Stop pretending to be confident, above hypocrisy." (V2/019). This data shows a violation of the maxim of agreement because the utterance does not open up space for agreement, but instead directly rejects and directs a negative assessment of the target party. The use of the imperative form "stop" signals an explicit rejection of an attitude perceived as "pretending to be self-confident," while the phrase "above hypocrisy" reinforces the value opposition between the speaker and the target. Thus, this utterance is not oriented toward dialogue or finding common ground, but rather directly emphasizes disagreement, thus violating the maxim of consensus.

2.3 Violation of the Sympathy Maxim

Sympathy was absent from several utterances included in the research data. The speaker demonstrated antipathy through utterances that justified or supported negative treatment of the target party. This condition reflects a lack of concern for the feelings and circumstances of the other party in the communication process. Violations of the sympathy maxim can be observed in the following data.

Data 6

"emg pantes d hujat" (V3/027). This data demonstrates a violation of the sympathy maxim because the utterance does not convey empathy but instead justifies the slanderous action against the target party. The expression "emg pantes d hujat" contains a judgment that justifies negative treatment without considering the humanity or condition of the person being spoken about. Thus, the comment reflects a lack of sympathy and a negative assessment that disrespects the target audience.



3. Classification of Hate Speech Based on Circular Letter from the Chief of Police Number SE/6/X/2015 and its Legal Implications under Law Number 1 of 2024 on Electronic Information and Transactions (ITE Law)

The classification of hate speech in this study refers to Circular Letter from the Chief of Police Number SE/6/X/2015 concerning Handling Hate Speech. The research data shows three forms of hate speech: insults, defamation, and provocation. Defamation, unpleasant acts, incitement, and the spread of false news were not found in the entire data. Each form of hate speech was then mapped based on its legal implications in accordance with the provisions of Law Number 1 of 2024 concerning Electronic Information and Transactions.

3.1 Insults

Insults were the most dominant form of hate speech in this study. This dominance indicates that speakers used words or phrases that demeaned the target's honor more frequently than other forms of hate speech. Insults were manifested through insults, ridicule, and negative labels that attacked personal aspects. One example of the insult category is presented below.

Data 7

"Why are you so rude? You're so complicated. You're just saying something." (V1/011). This data falls into the insult category based on the Chief of Police's Circular Letter No. SE/6/X/2015 because it contains derogatory remarks and uses harsh language toward the target. The phrases "Why are you so rude?", "You're so complicated," and "You're just saying something" indicate ridicule and negative labels that serve to demean the target's dignity. The utterance uses several Javanese words, namely "koen" (you), "eroh" (to know), "ngetuwel" (tangle), and "koyo" (like). The researcher uses

an asterisk (*) for the word "jem****" to maintain ethical scientific writing without losing the utterance's original meaning. Contextually, this form refers to the word "jembut," a vulgar term that literally refers to hair in the genital area. In this utterance, the word is not used in its denotative meaning, but rather as a form of insult to reinforce the expression of insult and demean the speaker. From an illocutionary perspective, this utterance is dominated by expressive speech acts indicating the speaker's dislike and may contain assertive elements in the form of negative assessments of others. Pragmatically, this data violates the maxim of appreciation because it maximizes the form of reproach. Therefore, this utterance has the potential to be highly associated with Article 27A of the ITE Law because it contains elements of degrading honor or reputation.

3.2 Defamation

Data in this category is characterized by speech containing specific accusations against an individual, negative claims presented as facts, and mentions of bad deeds without supporting evidence in the comments. In contrast to insults, which are more oriented towards labeling and expressive attacks on dignity, defamation emphasizes the construction of accusations aimed at damaging the reputation of the targeted party. Therefore, this category demonstrates a tendency to use assertive language to construct negative perceptions of a particular individual or group. The following is a description of the defamation category data.

Data 8

"After buying lots of like bombs, now buying lots of buzzers" (V2/014). This data falls into the defamation category based on the Chief of Police Circular Letter No. SE/6/X/2015 because it contains negative claims directed at a specific party, alleging unethical actions. The phrases "buying lots of like bombs" and "buying lots of buzzers" represent accusations of manipulative practices in the digital space, positioned as fact, despite the lack of evidence

in the context of the comments. From an illocutionary perspective, this utterance is dominated by assertive speech acts because it contains statements and claims constructed as truth by the speaker. Pragmatically, this data violates the maxim of appreciation because it maximizes accusations and negative assessments of the target party. Therefore, this utterance has a high level of indication of being related to Article 27A of the ITE Law because it has the potential to damage reputations through attributions of actions that harm the image of a particular party.

3.3 Provocation

Provocation differs from the previous two categories in that it is not only directed at the target of the utterance but also has the potential to influence other readers. Speech in this category triggers a negative response toward a particular party. This pattern demonstrates that the provocation in the data does not stand as a single opinion, but rather as an attempt to build a collective, confrontational response. The following is a description of the data for the provocation category.

Data 9

"emg pantes d hujat" (V3/027). This data falls into the provocation category based on the Chief of Police's Circular Letter No. SE/6/X/2015 because it contains speech that can strengthen the legitimacy of attacking or insulting a particular party. The phrase "pantes d hujat" has the potential to encourage other readers to respond similarly to the target. These characteristics categorize this data as a form of provocation based on the Chief of Police's Circular Letter No. SE/6/X/2015. The legal implications relate to Article 27A and can be linked to Article 28 paragraph (2) of Law Number 1 of 2024 concerning Information and Electronic Transactions if it meets the elements specified in the legislation.

DISCUSSION

Types of Illocutionary Speech Acts

Research findings indicate that hate speech in the YouTube comment section of the "Hilirisasi" series is determined not only by the choice of vocabulary, but also by the function of the speech act used in its delivery. Expressive speech acts were the most dominant form, followed by directive and assertive speech acts. This dominance indicates that speakers primarily use language as a means to express emotions, convey negative judgments, and influence readers' responses to the target audience. These characteristics demonstrate that the illocutionary function plays a significant role in shaping hate speech in the digital space.

This aligns with the views of Austin (1962:98) and Wijana (1996:18), who emphasized that illocutionary speech acts serve not only to express an opinion but also to carry out an action through speech. Based on data analysis, the dominant speech acts were expressive, assertive, and directive, as classified by Searle in Leech (2015:164). Expressive speech acts function to convey negative emotions such as anger, ridicule, and contempt toward a particular party. Assertive speech acts are used to convey claims or accusations that tend to have a negative connotation toward the subject being discussed. Meanwhile, directive speech acts appear in the form of provocation or encouragement that directs the response of other readers.

The findings of this study are similar to those of Surbakti et al. (2023), which showed that expressive and directive speech acts dominate hate speech on TikTok. This similarity indicates that hate speech in digital spaces tends to be constructed through emotional expressions and encouragement to others, even though the objects of study and the social media platforms used are different. The research of Adelawati et al. (2025) also shows that a pragmatic approach within a forensic linguistics framework is able to reveal the speaker's intent through speech acts that are not always conveyed literally. The difference in this study lies in the scope of the analysis, which not only identifies



the form of illocutionary speech acts but also links them to violations of politeness maxims, the classification of hate speech based on the Chief of Police Circular Letter Number SE/6/X/2015, and the legal implications according to Law Number 1 of 2024 concerning Information and Electronic Transactions.

These findings demonstrate that identifying hate speech is not solely based on the use of negative words; it is also necessary to consider the function of the speech acts that accompany each utterance. Illocutionary analysis provides insight into communicative intent, which is not always directly apparent in comments but still has the potential to attack honor or influence others. This approach strengthens the position of forensic linguistics as a study capable of explaining the relationship between language use, communication context, and the consequences that can arise in the digital space.

Types of Violations of Politeness Maxims

Language politeness is an aspect of pragmatics that plays a role in maintaining social relations between participants in a conversation. Rahardi (2005:35) states that the study of language politeness focuses on language use within a particular speech community, where politeness is manifested when speakers adhere to norms of politeness, thus maintaining harmonious social relations. In a pragmatic context, Leech (2015:206) explains that politeness maxims serve as guidelines for maintaining a balanced social relationship between speakers and their interlocutors, including in relations with third parties. However, in digital communication, not all speech adheres to these principles, resulting in violations of the maxim of politeness. Research findings indicate that this principle is frequently violated in the comments section of the "Hilirisasi" video series on the GibranTV YouTube account, leading to communication dominated by offensive speech rather than fostering polite interactions.

Based on data analysis of the comments section of the "Hilirisasi" YouTube series, the

dominant violations of the maxim of politeness were the maxim of appreciation, the maxim of sympathy, the maxim of wisdom, and the maxim of agreement. Violations of the maxim of appreciation are evident through the use of insults, ridicule, and negative labeling of certain parties. Violations of the maxim of sympathy indicate a lack of empathy for the target of the speech. Meanwhile, violations of the maxim of wisdom appear in the form of encouragement that has the potential to harm others, while violations of the maxim of agreement are evident in speech that reinforces conflict and polarizes opinion.

These research findings complement the research by Surbakti et al. (2023), which demonstrated the dominance of expressive and directive speech acts in hate speech on social media. This research focuses on the illocutionary function and therefore fails to explain how linguistic impoliteness contributes to the formation of hate speech. Research by Adelawati et al. (2025) also demonstrates that a pragmatic approach in forensic linguistics can reveal speakers' intentions through non-literal speech acts containing irony, sarcasm, hyperbole, and humor. This study expands on these findings by demonstrating that verbal aggression is manifested not only through the form of speech but also through violations of the principles of politeness that govern communication ethics. The results of research by Farah et al. (2025) also show that hate speech in social media comment sections is dominated by insults manifested through the choice of harsh lexicon and attacks on the target's personality. This study expands on these findings by demonstrating that the use of negatively valued lexicon not only reflects the form of hate speech but also manifests as a violation of the maxims of appreciation, agreement, and sympathy, which are indicators of verbal aggression in digital communication.

The difference between this study and previous research lies in the integration of the analysis of illocutionary speech acts and violations of politeness maxims within a single forensic linguistic framework. This approach allows for a

more comprehensive identification of hate speech because it considers both the function of language and the quality of interpersonal relationships built through speech. This study also uses comment data on public policy issues on YouTube and refers to the latest regulation, namely Law Number 1 of 2024 concerning Electronic Information and Transactions, thus providing a more up-to-date perspective compared to previous studies that generally use Instagram or TikTok data with reference to the 2016 ITE Law. These findings indicate that violations of politeness maxims are not only indicators of linguistic impoliteness, but can also be an early indicator of the emergence of verbal aggression that has the potential to develop into hate speech in the digital space.

Classification of Hate Speech Based on Circular Letter of the Chief of Police Number SE/6/X/2015 and its Legal Implications under the Electronic Information and Transactions Law Number 1 of 2024 (ITE Law)

In forensic linguistics, language is understood not only as a means of communication but also as linguistic evidence that can be systematically analyzed in a legal context. Coulthard & Johnson (2007) explain that forensic linguistics focuses on the analysis of language in various legal contexts to uncover the meaning, intent, and potential implications of an utterance. In this view, every utterance can be treated as linguistic data with interpretive value in relation to legal issues.

The forensic linguistic construct is built by viewing comments on YouTube columns as a form of linguistic data that can be analyzed to identify tendencies toward hate speech based on linguistic indicators. The analysis is not aimed at establishing legal guilt, but rather at demonstrating how language can function as linguistic evidence related to Circular Letter of the Chief of Police Number SE/6/X/2015 and the ITE Law Number 1/2024. Thus, forensic linguistics serves as a bridge between language analysis and linguistic mapping of legal categories.

Research findings indicate that the comment section of the "Downstreaming" video series on the GibranTV YouTube account is not only used as a medium for expressing opinions on public policy, but also as a space for the emergence of hate speech through insults, defamation, and provocation. This situation aligns with the views of Mintowati, Nasrullah, & Ahmadi (2025:39), who state that hate speech is speech used to attack specific individuals or groups, and the view of Sari, Pastika, and Satyawati (2025) that hate speech has the potential to cause prejudice, discrimination, and social conflict. In this study, the dominant forms of hate speech were insults, defamation, and provocation directed at the target of the comments. The dominance of these three forms of speech indicates that criticism of policies in the digital space can shift into personal attacks containing elements of hatred.

The emergence of these forms of speech cannot be separated from the characteristics of digital communication, which provides space for users to express their opinions freely. Andzani & Irwansyah (2023:1967) explain that the dynamics of digital communication are influenced by a participatory culture and the high level of netizen involvement in responding to emerging issues in cyberspace. Furthermore, the relative anonymity of digital spaces allows users to express their opinions more freely and directly and aggressively. This is evident in comments that no longer focus on the substance of downstream policies, but instead shift to personal attacks through the use of insults, accusations, and negative labeling. This finding also supports the opinion of Mintowati, Nasrullah, & Ahmadi (2025:74) that hate speech in digital media is often disguised through satire, irony, or other forms of expression that appear as ordinary criticism.

The research findings have legal implications if they meet the elements stipulated in Law Number 1 of 2024 concerning Electronic Information and Transactions. Speech that attacks a person's honor or reputation through electronic media has the



potential to meet the elements of Article 27A of the ITE Law Number 1 of 2024, while provocative speech can be considered a legal offense if it meets the applicable elements of a criminal offense. This analysis is not intended to establish a violation of the law, but rather to demonstrate that linguistic indicators are related to legal provisions and can therefore serve as linguistic evidence in forensic studies.

These findings align with research by Nurlisma (2022), Farah et al. (2025), and Komara et al. (2025), which all found a predominance of insults and provocation in hate speech on social media. The difference is that these studies still refer to the 2016 Electronic Information and Transactions (ITE) Law and have not integrated analysis of illocutionary speech acts and violations of politeness maxims. This study offers a more comprehensive approach by connecting illocutionary functions, violations of politeness maxims, the classification of hate speech based on the Chief of Police's Circular Letter No. SE/6/X/2015, and the legal implications based on the ITE Law No. 1 of 2024 within a single forensic linguistics framework. This approach represents a novel research approach in examining hate speech on social media, particularly in the context of public policy discourse on YouTube.

CONCLUSION

Based on the results of research into hate speech in the comments section of the YouTube series "Hilirisasi" on the @GibranTV channel, through a study of illocutionary speech acts, violations of politeness maxims, and forensic linguistics, it can be concluded that speech in the comments section functions not only as an expression of opinion but also as a form of verbal response that carries the potential for linguistic aggression. The identified hate speech appears not only in the form of direct insults but also in the form of accusations, negative labels, invitations,

and provocations constructed through specific linguistic strategies.

Within the framework of illocutionary speech acts, hate speech is predominantly realized through expressive, assertive, and directive speech acts, each of which serves to express negative emotions, convey claims or accusations, and direct or influence the actions of others. This demonstrates that the meaning of hate speech can be seen not only from its lexical form but also from the communicative function underlying it in digital interactions.

From a linguistic politeness perspective, violations of the maxims of praise, sympathy, wisdom, agreement, generosity, and simplicity were found, indicating that netizens' communication tended not to maintain harmonious interactions and was more oriented toward verbal attacks against certain parties. Meanwhile, from a forensic linguistics perspective, these utterances can be identified as forms of insult, provocation, incitement, defamation, and indications of the dissemination of unverified information based on the Chief of Police Circular Letter Number SE/6/X/2015 and the Electronic Information and Transactions Law Number 1 of 2024. Thus, language in YouTube comment sections is not only a means of communication but also has the potential to be a medium for verbal aggression that can be analyzed pragmatically and forensically in the context of digital communication.

ACKNOWLEDGMENTS

The author would like to express his gratitude to all parties who have contributed to the development of this research. He would like to express his gratitude to his supervisor, his beloved family, his supportive colleagues, himself, and others who have assisted in this research process. I also thank netizens who have commented on the video series themed "Downstreaming" on the YouTube account @GibranTV.

REFERENCES

- Adelawati, M & Arimi, S. (2025). Tindak Tutur Tidak Literal pada Pelanggaran UU ITE dalam Kasus Lina Mukherjee: Analisis Linguistik Forensik. *Semantik*, 14(1), 113–126. <https://doi.org/10.22460/semantik.v14i1.p113-126>
- Andzani, D., & Irwansyah, &. (2023). Dinamika Komunikasi Digital: Tren, Tantangan, dan Prospek Masa Depan. *JURNAL SYNTAX DMIRATION*, 4(11), 1964–1976. <https://doi.org/hIps://doi.org/10.46799/jsa.v4i7.671> DINAMIKA
- Austin, J. L. (1962). *How to do Things with Words*. Oxford : Oxford University Press.
- Coulthard, M & Johnson, A. (2007). *An Introduction to Forensic Linguistics: Language in Evidence*. London: Routledge.
- Ernawati & Prayitno, H. J. (2024). Penggunaan dan Penyimpangan Prinsip Kesantunan Berbahasa pada Komentar Kontra Warganet Dikanal YouTube Gita Safitri Devi dalam Konten Video “*Childfree*: Serba Salah Dimata Warganet.” *J-Symbol: Jurnal Magister Pendidikan Bahasa dan Sastra Indonesia*, 12(1), 70—84. <https://doi.org/https://doi.org/10.23960/J-Symbol>
- Farah, N. A., Salimah, H., Nurul, I., Zaimah, S., & Isnaini, L. (2025). Analisis Linguistik Forensik Ujaran Kebencian Warganet pada Kolom Komentar Akun Instagram @aldymldni dalam Kasus Penipuan yang Dilakukan oleh Aldy Maldini. *@Artikulasi Jurnal Pendidikan Bahasa dan Sastra Indonesia*, 5(2), 139–161. <https://doi.org/https://doi.org/10.17509/artikulasi.v5i2.86229>
- Komara, A. B. Hasani, A. & Firmansyah, D. (2025). Analisis Linguistik Forensik terhadap Ujaran Kebencian Netizen pada Kolom Komentar Instagram Kaesang Pangarep dan Erina Gudono Tahun 2024. *Pendas: Jurnal Ilmiah Pendidikan Dasar*, 10(03), 200–212. <https://doi.org/https://doi.org/10.23969/jp.v10i3.28328>
- Leech, G. (2015). *Prinsip-Prinsip Pragmatik* (1st ed.). Jakarta: Universitas Indonesia.
- McMenamin, G. R. (2002). *Forensic Linguistics: Advances in Forensic Stylistics*. Boca Raton, Florida: CRC Press.
- Mintowati, M., Nasrullah, N., & Ahmadi, A. (2025). *Linguistik Forensik di Era Digital*. Lampung: Literasi Nusantara Abadi Grup.
- Nurlisma, N. (2022). *Ujaran Kebencian terhadap Artis Nissa Sabyan di Media Sosial (Kajian Linguistik Forensik)*. Tarakan: Universitas Borneo.
- Rahardi, K. (2005). *Pragmatik: Kesantunan Imperatif Bahasa Indonesia*. Jakarta: Erlangga.
- Razak, A. (2017). *Metode Riset: Menggapai Mixed Methods Bidang Pembelajaran Bahasa Indonesia*. Pekanbaru: Ababil Press.
- Sari, M. P. D., Pastika, I. W., & Satyawati, M. S. (2025). Ujaran Kebencian terhadap Selebgram Azizah Salsha di Media Sosial TikTok: Kajian Linguistik Forensik. *Kulturistik: Jurnal Ilmu Bahasa dan Budaya*, 9(2), 49–57.
- Surbakti, D. Br., Sihombing, D. N., & Nadira, A. J. (2023). Linguistik Forensik Ujaran Kebencian terhadap Artis Tiara Andini di Media Sosial TikTok. *Didaktik: Jurnal Ilmiah PGSD FKIP Universitas Mandiri*, 09(05), 2337–2344.
- Syafruddin, S., Thaba, A. & Ananda, R. (2024). Ujaran Kebencian Netizen Indonesia pada Akun Twitter Es Teh: Tinjauan Linguistik Forensik. *Semantik*, 13(1), 15–28.
- Wijana, I. D. P. (1996). *Dasar-Dasar Pragmatik*. Yogyakarta: Andi Offset.
- Yuliani, W. (2018). Metode Penelitian Deskriptif Kualitatif dalam Perspektif Bimbingan dan Konseling. *QUANTA: Jurnal Kajian Bimbingan dan Konseling dalam Pendidikan*, 2 (2), 83–91.