

AI Threat or Partner Narratives in Lecturer Student Discourse: A CDA of Polarization and Legitimacy

Rama Dwika Herdiawan¹, Yayah Nurhidayah², Trian Pamungkas³,
Yaniar Ratna Dewi⁴

^{1,2,3,4} English Language Education Department, Faculty of Teacher Training and Education, Universitas Majalengka, Indonesia

Corresponding Author: ramadwika@unma.ac.id

Article History: Submitted date February 14th, 2026; Revised date March 12th, 2026; Accepted date June 4th, 2026; Published date June 30th, 2026

ABSTRACT

This study examines how lecturers and students discursively construct generative artificial intelligence (GenAI) as either a threat or a partner in higher education, and how these competing stances are legitimized in everyday academic communication. Guided by Critical Discourse Analysis, the study analyzes 124 GenAI-relevant statements (74 spoken; 50 written) drawn from assignment briefings, classroom discussions, consultation meetings, LMS announcements, rubric explanations, emails, and class-group chats in a bounded university context. Polarization was coded as Threat, Partner, or Hybrid/Ambivalent, while legitimation strategies were traced using Theo van Leeuwen's framework (authorization, moral evaluation, rationalization, and mythopoesis), complemented by discursive legitimation perspectives. Findings show that Threat-oriented statements (n=52) slightly outnumber Partner-oriented statements (n=35), with Hybrid "allowed, but..." regulation (n=37) functioning as the dominant mechanism for stabilizing norms under uncertainty. Threat talk occurs during high-stakes assessment moments and is commonly legitimized through authorization (rules, rubrics, sanctions) and moral evaluation (fairness, honesty), often reinforced by cautionary narratives. Partner talk becomes more salient in formative learning episodes and is primarily legitimized through rationalization that frames GenAI as scaffolding for brainstorming, outlining, summarizing, and language polishing, while reasserting student responsibility for accuracy and voice. Across stances, the central legitimacy struggle concerns accountability, who remains responsible for claims, originality, and transparency when GenAI is involved alongside student anxiety produced by inconsistent boundaries across courses. The study argues that sustainable governance requires clearer communicative templates, process-oriented assessment evidence (draft trails, reflective logs, brief oral explanations), and

How to Cite (in APA 7th Edition):

Herdiawan, R. D., Nurhidayah, Y., Pamungkas, T., & Dewi, Y. R. (2026). AI Threat or Partner Narratives in Lecturer Student Discourse: A CDA of Polarization and Legitimacy. *Lensa: Kajian Kebahasaan, Kesusastraan, Dan Budaya*, 16(1), 79–102. <https://doi.org/10.26714/lensa.16.1.2026.79-102>

careful use of detection tools only as triggers for follow-up review to protect fairness and trust.

Keywords: GenAI; integrity; discourse; legitimation; interaction; governance; ambiguity

INTRODUCTION

The rapid diffusion of generative artificial intelligence (GenAI), especially large language models integrated into conversational interfaces such as ChatGPT, has reconfigured how academic work is produced, evaluated, and justified in higher education (Dwivedi et al., 2023; Kasneci et al., 2023). In only a short time, GenAI has moved from a "novel tool" to an everyday infrastructure for drafting, paraphrasing, summarizing, and ideating, creating a new normal in which boundaries between human authorship, machine assistance, and collaborative writing practices are no longer self-evident (Dwivedi et al., 2023). Education-focused syntheses underline that large language models bring simultaneous pedagogical promise (e.g., feedback, scaffolding, personalization) and non-trivial risks (e.g., hallucination, opacity, bias, dependency), making them uniquely disruptive for assessment cultures that still rely heavily on textual products as evidence of learning (Kasneci et al., 2023).

In this context, "academic integrity" has become one of the most visible flashpoints. Empirical and conceptual work shows that GenAI intensifies long-standing tensions in higher education assessment, particularly the difficulty of distinguishing legitimate support (peer feedback, proofreading, digital tools) from practices interpreted as misrepresentation (Cotton et al., 2023; Luo, 2024). As institutions respond, integrity is increasingly framed not only as a student moral obligation, but also as a governance problem that requires clearer norms, shared language, and enforceable boundaries (Cotton et al., 2023). Cotton et al. (2023), for example, argue that the emergence of ChatGPT compels universities to rethink how integrity is defined and operationalized in assessment practice, as new forms of machine-mediated writing contest traditional markers of originality and independent authorship.

Alongside integrity concerns, a second shift is occurring: universities are producing new policy texts and guidance documents that attempt to stabilize meaning and practice around GenAI. Policy analyses suggest that institutional responses are not merely technical (what tools are allowed), but also ethical and political (who is responsible, what counts as transparency, and which values are prioritized) (Dabis & Csáki, 2024; Wang et al., 2024). Dabis and Csáki's (2024) analysis of early policy responses emphasizes that universities' first frameworks repeatedly foreground ethical dimensions such

How to Cite (in APA 7th Edition):

erdiawan, R. D., Nurhidayah, Y., Pamungkas, T., & Dewi, Y. R. (2026). AI Threat or Partner Narratives in Lecturer Student Discourse: A CDA of Polarization and Legitimacy. *Lensa: Kajian Kebahasaan, Kesusastraan, Dan Budaya*, 16(1), 79–102. <https://doi.org/10.26714/lensa.16.1.2026.79-102>

as accountability, transparency, inclusiveness, and human oversight, revealing that GenAI is being governed as a legitimacy-sensitive change rather than a neutral technology adoption.

However, the governance challenge cannot be separated from the longer history of digital writing. Writing in higher education has long been supported by layered tools search engines, grammar checkers, citation managers, translation software, and collaborative platforms. [Zhao et al. \(2024\)](#) caution that GenAI use should be interpreted within this broader "digitisation of writing," because students already assemble tool ecologies to manage writing as an iterative and cognitively demanding process. In this sense, the controversy around ChatGPT is not simply about a new tool; it is also about which forms of tool-supported writing can be publicly justified as legitimate academic practice ([Zhao et al., 2024](#)).

These structural shifts have also produced a striking polarization in how GenAI is discussed across educational communities: AI is alternately constructed as a threat (to integrity, learning authenticity, fairness, and disciplinary standards) and as a partner (for learning support, accessibility, creativity, and efficiency). Evidence from cross-cultural perspectives indicates that stakeholders often interpret GenAI through this paradoxical lens, debating whether it is primarily a threat to academic integrity or a catalyst for reform in assessment and pedagogy ([Yusuf et al., 2024](#)). This polarization is not merely rhetorical noise; it actively shapes classroom expectations, student risk-taking, and institutional policy trajectories ([Yusuf et al., 2024](#)).

At the level of everyday educational experience, the polarization becomes visible in how lecturers and students explain, contest, and normalize the use of GenAI. Large-scale student evidence indicates that learners often value ChatGPT for practical academic purposes (e.g., brainstorming, summarization, simplifying complexity), while simultaneously expressing concern about cheating and supporting calls for regulation ([Ravšelj et al., 2025](#)). Such findings point to an ambivalent but consequential interpretive environment: students may adopt GenAI as a functional learning partner while anticipating that lecturers or institutions may interpret similar practices as integrity violations ([Ravšelj et al., 2025](#)).

Crucially, this ambivalence is mediated by discourse especially the micro-level statements exchanged in classrooms, supervision meetings, assignment briefings, and informal academic conversations. Yet much of the GenAI-in-education literature has prioritized (a) perception surveys, (b) policy and guideline analysis, or (c) technical implications for assessment design, leaving the interactional discourse of lecturer–student meaning-making comparatively underexplored ([Wang et al., 2024](#)). Policy-focused research demonstrates that universities frame GenAI through particular categories and intervention areas, shaping what counts as ethical or effective use; for instance, analyses of institutional resources show how universities define acceptable practices and distribute responsibility through guidance structures ([Wang et al., 2024](#)). At the same time, teacher-focused work

How to Cite (in APA 7th Edition):

erdiawan, R. D., Nurhidayah, Y., Pamungkas, T., & Dewi, Y. R. (2026). AI Threat or Partner Narratives in Lecturer Student Discourse: A CDA of Polarization and Legitimacy. *Lensa: Kajian Kebahasaan, Kesusastraan, Dan Budaya*, 16(1), 79–102. <https://doi.org/10.26714/lensa.16.1.2026.79-102>

indicates that readiness and prior beliefs influence educators' intentions and can generate dichotomous stances toward ChatGPT sometimes even positioning the tool as a competitor, suggesting that polarized discourse is likely to surface in pedagogical talk, not only in formal policy texts ([Rahimi & Sevilla-Pavón, 2024](#)).

This study situates polarization as a discursive struggle over legitimacy enacted in lecturer–student statements. In this paper, GenAI refers to AI systems capable of generating human-like language (and other content) from prompts; polarization refers to the discursive production of binary oppositions (here, "AI as a Threat" vs. "AI as a Partner") that simplify a complex sociotechnical change into competing moral and practical stances. Legitimacy/legitimation refers to how discourse makes certain actions, identities, and norms appear acceptable, reasonable, and "proper" within an institution. From a critical discourse analysis (CDA) perspective, the focus shifts from whether GenAI is objectively good or bad to how educational actors construct the technology, assign responsibility, and normalize or stigmatize practices through language. In particular, Theo van Leeuwen's legitimation framework conceptualizes legitimacy as something discourse does, for example, through appeals to authority and rules, moral evaluation, rationalization, and narrative forms that warn or reward ([van Leeuwen, 2007](#)).

From a Critical Discourse Analysis (CDA) perspective, the debate on GenAI in higher education is not merely a disagreement about technology, but a struggle over authority, responsibility, and legitimate academic practice. Drawing on Fairclough's CDA tradition, this study views lecturer–student discourse as part of a broader institutional process through which values such as originality, honesty, transparency, and accountability are negotiated. Fairclough's approach provides the critical basis for linking language to power and institutional practice, while van Leeuwen's legitimation framework offers specific tools for identifying how certain positions are justified in discourse. Through authorization, moral evaluation, rationalization, and mythopoesis, discourse can make particular practices appear acceptable, ethical, reasonable, or risky ([van Leeuwen, 2007](#)). In this study, the framework is not applied only as a descriptive coding tool. Instead, it explains how the polarization of GenAI as either a "threat" or a "partner" serves as a form of legitimation. Lecturer–student statements therefore do not simply describe GenAI use; they actively shape what counts as acceptable assistance, authentic authorship, and responsible academic conduct. For this reason, the study is positioned as a critical CDA inquiry rather than merely an applied discourse analysis, as it examines how everyday academic discourse produces, contests, and stabilizes institutional meanings around GenAI use.

Discourse-oriented analyses of university GenAI guidance further underline that institutional texts and everyday academic talk can frame agency, accountability, and inclusivity in ways that shape how "proper" GenAI use becomes speakable and enforceable ([Simpson, 2025](#)). This study therefore treats "AI as a Threat vs. Partner" as a discursive struggle over

How to Cite (in APA 7th Edition):

erdiawan, R. D., Nurhidayah, Y., Pamungkas, T., & Dewi, Y. R. (2026). AI Threat or Partner Narratives in Lecturer Student Discourse: A CDA of Polarization and Legitimacy. *Lensa: Kajian Kebahasaan, Kesusastraan, Dan Budaya*, 16(1), 79–102. <https://doi.org/10.26714/lensa.16.1.2026.79-102>

legitimacy. From a critical discourse analysis (CDA) perspective, the issue is not simply whether GenAI is good or bad, but how particular statements authorize, moralize, or rationalize positions that make certain practices appear normal and others deviant. This study therefore, treats "AI as a Threat vs. Partner" as a discursive struggle over legitimacy. From a Critical Discourse Analysis (CDA) perspective, the central issue is not whether GenAI is inherently good or bad, but how discourse makes particular academic practices appear acceptable, ethical, risky, or deviant. Existing CDA and discourse studies on GenAI have largely examined institutional policies, guidelines, and public debates, thereby showing how universities formally frame responsibility, transparency, and ethical use. However, such studies have paid less attention to everyday lecturer–student discourse, where these institutional meanings are interpreted, negotiated, and sometimes contested in actual academic interaction. This gap is significant because GenAI governance is not enacted only through official documents; it is also reproduced through assignment briefings, classroom warnings, consultation talk, LMS announcements, and informal exchanges. By analyzing these micro-level statements, this study extends CDA work on GenAI from policy-level representation to interactional legitimation. It shows that polarization itself functions as a legitimating strategy: the construction of GenAI as a "threat" legitimizes restriction, suspicion, and control, while its construction as a "partner" legitimizes scaffolding, assistance, and conditional acceptance. Thus, the study offers a critical CDA account of how academic authority, responsibility, and student identity are discursively produced in the everyday governance of GenAI use.

Recent discourse-oriented work strengthens the need for this focus. [Simpson's \(2025\)](#) critical discourse analysis of university GenAI guidance shows that institutional texts do not merely inform; they frame agency, accountability, and acceptable academic identities in ways that can constrain how lecturers and students talk about "proper" use. Meanwhile, assessment scholarship argues that policy responses frequently revert to narrow notions of originality, even as GenAI makes "originality" harder to operationalize in defensible ways, prompting calls to reconsider how student work is interpreted and assessed ([Luo, 2024](#)).

Accordingly, the present research investigates how lecturer–student statements construct and contest the polarity of "AI as a threat vs. partner," and how these constructions are legitimized (or delegitimized) through discourse. By analyzing polarization and legitimation in situated academic talk, this study aims to contribute (1) an empirically grounded account of how GenAI norms are negotiated at the micro level, (2) a CDA-informed explanation of the rhetorical strategies that naturalize certain stances, and (3) implications for policy communication and assessment design that are sensitive to lived classroom discourse rather than policy ideals alone ([Luo, 2024](#); [Wang et al., 2024](#)).

How to Cite (in APA 7th Edition):

erdiawan, R. D., Nurhidayah, Y., Pamungkas, T., & Dewi, Y. R. (2026). AI Threat or Partner Narratives in Lecturer Student Discourse: A CDA of Polarization and Legitimacy. *Lensa: Kajian Kebahasaan, Kesusastraan, Dan Budaya*, 16(1), 79–102. <https://doi.org/10.26714/lensa.16.1.2026.79-102>

METHOD

Research Design

This study employs a qualitative design grounded in Critical Discourse Analysis (CDA) to investigate how lecturer–student statements construct the polarization of "AI as a Threat vs. Partner" and how these positions are discursively legitimized or delegitimized in everyday academic communication. CDA is suitable because the GenAI controversy in higher education is not merely technical or pedagogical; it is a contestation over norms, responsibility, and acceptable academic identities that becomes visible through language choices in warnings, permissions, advice, and justifications (Simpson, 2025; Wang et al., 2024; Herdiawan, R. D., & Windiarti, D., 2025). The analysis is anchored in legitimation as a discourse function, using van Leeuwen's categories—authorization, moral evaluation, rationalization, and mythopoesis to trace how lecturers and students justify restrictions or endorsements of GenAI use through appeals to rules, values, practical reasoning, and cautionary narratives (van Leeuwen, 2007; Herdiawan, R. D., Fakhruddin, A., Nurhidayat, E., & Nurhadi, K., 2025). To strengthen the CDA lens on legitimation as strategic discourse, the study also draws on discursive legitimation perspectives that emphasize how legitimacy is built through rhetorical and framing moves within institutional settings (Vaara & Tienari, 2008).

Research Setting, Participants, and Sampling Strategy

This study was situated in a bounded Indonesian higher education context, namely a private university in Majalengka Regional Regency, within English Language Education courses that required writing- and assessment-based tasks. This context was selected because GenAI debates become particularly visible when written products are used as evidence of learning, originality, and academic integrity (Cotton et al., 2023; Luo, 2024). Data were collected over 4 months from 8 lecturers in the Bachelor of English Language Education and 88 undergraduate students who participated in GenAI-related academic communication across classroom interaction, assignment briefings, consultation meetings, LMS announcements, emails, and class-group chats. Participants were recruited through purposive sampling because they were directly involved in producing or responding to discourse about the use of GenAI in academic tasks. To strengthen the variation of the corpus, maximum variation sampling was applied using explicit criteria. Lecturer participants were selected based on three considerations: their involvement in writing- or assessment-oriented courses, their role in communicating rules or expectations about GenAI use, and their different orientations toward GenAI, ranging from restrictive to conditional and supportive. Student participants were selected based on their enrollment in the selected courses,

How to Cite (in APA 7th Edition):

herdiawan, R. D., Nurhidayah, Y., Pamungkas, T., & Dewi, Y. R. (2026). AI Threat or Partner Narratives in Lecturer Student Discourse: A CDA of Polarization and Legitimacy. *Lensa: Kajian Kebahasaan, Kesusastraan, Dan Budaya*, 16(1), 79–102. <https://doi.org/10.26714/lensa.16.1.2026.79-102>

their exposure to GenAI-related instructions, their academic achievement levels, and their reported degrees of GenAI engagement. The final sample was justified by the depth and relevance of the resulting discourse corpus rather than by statistical representativeness. The corpus consisted of 124 GenAI-relevant statements, which were sufficient to capture repeated patterns of threat, partner, and hybrid/ambivalent framings across spoken and written academic communication. This sampling logic is appropriate for a qualitative CDA study because the main concern is not the numerical distribution of participants, but the richness of discourse through which authority, responsibility, and legitimacy are negotiated in a specific institutional setting. This sampling logic is aligned with evidence that many higher-education students perceive GenAI tools as useful for academic tasks while simultaneously expressing integrity concerns, creating fertile ground for "threat/partner" polarization in classroom discourse (Ravšelj et al., 2025). In addition, educational research has highlighted that large language models offer both learning affordances and serious risks (e.g., hallucinations, bias, overreliance), which can shape divergent lecturer and student stances and thereby enrich the discourse corpus (Kasneci et al., 2023).

Data Sources and Corpus Construction

The dataset consists of lecturer–student statements gathered from both spoken interaction and written academic communication, enabling analysis of how legitimacy is produced across settings and genres of academic life. Spoken data may include audio-recorded classroom episodes (e.g., assignment briefings, integrity reminders, and explicit discussions of AI use), consultation or supervision meetings, and brief advising exchanges in which norms are negotiated in real time. Written data may include LMS announcements, assignment sheets, rubric explanations, email messages, and class-group chat excerpts when they contain explicit or implicit references to GenAI (e.g., permitted uses, restrictions, disclosure rules, or warnings about detection). The corpus is supplemented with institutional GenAI guidance documents—where available—as intertextual context, because policy and guidance texts often pre-structure how responsibility, transparency, and "acceptable use" are framed and then echoed (or resisted) in everyday talk (Dabis & Csáki, 2024; Wang et al., 2024). This intertextual approach also aligns with the argument that GenAI should be understood within the broader trajectory of digitized academic writing practices and tool ecologies, which influences what can be defended as legitimate support (Zhao et al., 2024).

How to Cite (in APA 7th Edition):

erdiawan, R. D., Nurhidayah, Y., Pamungkas, T., & Dewi, Y. R. (2026). AI Threat or Partner Narratives in Lecturer Student Discourse: A CDA of Polarization and Legitimacy. *Lensa: Kajian Kebahasaan, Kesusastraan, Dan Budaya*, 16(1), 79–102. <https://doi.org/10.26714/lensa.16.1.2026.79-102>

Operational Definitions and Unit of Analysis

For analytic precision, the primary unit of analysis is a statement, defined as a bounded spoken turn or written segment that performs at least one of the following functions: (a) evaluates GenAI (e.g., beneficial/harmful), (b) instructs, restricts, or permits its use (allowed/forbidden/conditional), (c) justifies a stance (why permitted/why risky), or (d) attributes responsibility (who is accountable for accuracy, originality, or disclosure). Polarization is operationalized as discourse that constructs a binary or near-binary opposition, positioning GenAI as a threat (e.g., cheating, dependence, unfairness, loss of authenticity) versus as a partner (e.g., scaffolding, efficiency, accessibility, ideation). Legitimacy/legitimation is operationalized as the discursive process through which speakers/writers make one side of the polarity appear reasonable, moral, rule-consistent, or inevitable through recognizable strategies, especially those mapped by van Leeuwen's framework ([van Leeuwen, 2007](#)), with additional sensitivity to legitimation as strategic framing in institutional discourse ([Vaara & Tienari, 2008](#)).

Data Collection Procedures and Ethics

Data collection proceeds in staged, ethics-forward procedures. First, focal courses are identified and participants are invited through informed consent protocols that explicitly acknowledge power relations (lecturer–student) and emphasize that participation is voluntary and non-evaluative (i.e., unrelated to grading). Second, naturally occurring interactions are recorded at agreed moments (e.g., assessment briefings), and written materials are collected only with participants' consent. Third, all recordings are transcribed verbatim, anonymized through pseudonyms and removal of identifying details, and stored securely. Because GenAI governance is a legitimacy-sensitive domain closely tied to integrity, originality, and potential accusations of misconduct, confidentiality is treated as central; excerpts used in reporting are minimally traceable and edited only to remove identifiers without altering meaning ([Cotton et al., 2023](#); [Luo, 2024](#)). This ethical posture is consistent with the observation that institutional responses to GenAI often foreground accountability and transparency, which can shape what participants feel safe to say and how they perform legitimacy in talk ([Dabis & Csáki, 2024](#)).

Analytical Framework and Coding Scheme

Analysis follows an iterative CDA procedure that integrates (1) polarity mapping and (2) legitimation analysis. First, AI-relevant statements are identified through keyword scanning (e.g., "AI," "ChatGPT," "paraphrase tool," "detector," "declare/disclose") and close reading, then coded for stance orientation: Threat, Partner, or Hybrid/Ambivalent (including

How to Cite (in APA 7th Edition):

erdiawan, R. D., Nurhidayah, Y., Pamungkas, T., & Dewi, Y. R. (2026). AI Threat or Partner Narratives in Lecturer Student Discourse: A CDA of Polarization and Legitimacy. *Lensa: Kajian Kebahasaan, Kesusastraan, Dan Budaya*, 16(1), 79–102. <https://doi.org/10.26714/lensa.16.1.2026.79-102>

conditional permissions such as allowing brainstorming but prohibiting drafting). Second, each statement is coded for legitimation strategies using van Leeuwen's categories: authorization, moral evaluation, rationalization, and mythopoesis—to capture how rules, values, instrumental reasoning, and narratives are mobilized to justify or contest GenAI use ([van Leeuwen, 2007](#)). This step is strengthened by insights from discursive legitimation that highlight how institutional actors use rhetorical strategies to stabilize contested practices and identities ([Vaara & Tienari, 2008](#)). After this first-cycle coding, a second-cycle synthesis consolidates recurring patterns and themes across contexts, using a thematic consolidation logic to organize how particular legitimation strategies cluster around "threat" versus "partner" framings ([Braun & Clarke, 2006](#)).

Analytic Procedure, Rigor, and Trustworthiness

To enhance rigor and trustworthiness, the study maintains an explicit audit trail of analytic decisions, coding revisions, and memo-based interpretations, consistent with qualitative trustworthiness guidance ([Nowell et al., 2017](#)). A portion of the dataset is double-coded to test the stability of category boundaries (e.g., distinguishing moral evaluation from rationalization or identifying hybrid stances), and disagreements are treated as analytic opportunities to refine the codebook rather than as purely technical errors, aligning with practical recommendations for intercoder processes in qualitative research ([O'Connor & Joffe, 2020](#)). The final interpretive stage links micro-level lecturer–student discourse to the wider controversy documented in recent literature, particularly integrity anxieties and contested originality norms ([Cotton et al., 2023](#); [Luo, 2024](#)), the governance role of institutional guidance ([Dabis & Csáki, 2024](#); [Wang et al., 2024](#)), and the broader digitization of writing that reshapes what counts as legitimate academic support ([Zhao et al., 2024](#)). Through this linkage, the study explains not only what polarized positions look like in interaction but also how legitimacy is discursively produced and circulated between policy-oriented texts and everyday pedagogical talk ([Simpson, 2025](#)).

FINDINGS AND DISCUSSION

Findings

Across the corpus, the analysis examines how lecturer–student statements construct a polarized discourse of "AI as a Threat vs. Partner" and how each position is legitimized or delegitimized in everyday academic communication. The 124 GenAI-relevant statements were produced by [number] participants, consisting of [number] lecturers and [number] undergraduate students from [discipline/program] at [type/name of institution] in Indonesia. Of these statements, [number] were produced by lecturers and [number] by students.

How to Cite (in APA 7th Edition):

erdiawan, R. D., Nurhidayah, Y., Pamungkas, T., & Dewi, Y. R. (2026). AI Threat or Partner Narratives in Lecturer Student Discourse: A CDA of Polarization and Legitimacy. *Lensa: Kajian Kebahasaan, Kesusastraan, Dan Budaya*, 16(1), 79–102. <https://doi.org/10.26714/lensa.16.1.2026.79-102>

The corpus was drawn from both spoken and written academic communication to capture how GenAI norms were negotiated across different interactional settings. Spoken data consisted of 74 statements from assignment briefings [28], classroom discussions [26], and consultation meetings [20]. Written data consisted of 50 statements from LMS announcements [15], rubric explanations [10], emails [12], and class-group chats [13].

Table 1.
Data Distribution

Data source	Type	Number of statements	Main producer
Assignment briefings	Spoken	28	Lecturers
Classroom discussions	Spoken	26	Lecturers and students
Consultation meetings	Spoken	20	Lecturers and students
LMS announcements	Written	15	Lecturers
Rubric explanations	Written	10	Lecturers
Emails	Written	12	Lecturers and students
Class-group chats	Written	13	Students and lecturers
Total		124	

The stance coding resulted in three categories: Threat-oriented statements ($n = 52$), Partner-oriented statements ($n = 35$), and Hybrid/Ambivalent statements ($n = 37$). These frequencies are not used to claim statistical generalization. Rather, they function as an interpretive map for identifying which discourse patterns became more visible and recurrent in the corpus. The higher number of Threat-oriented statements is analytically significant because it shows that GenAI was more frequently discussed in terms of risk, control, academic integrity, and assessment fairness than in terms of learning support. However, the substantial number of Hybrid/Ambivalent statements indicates that the discourse was not organized as a simple binary opposition. Instead, many statements combined permission with restriction, typically through conditional formulations such as "allowed, but...", "only for...", or "as long as...." This pattern suggests that polarization was stabilized through negotiated boundaries: GenAI could be accepted as a learning partner only when its use was made accountable through disclosure, verification, and

How to Cite (in APA 7th Edition):

erdiawan, R. D., Nurhidayah, Y., Pamungkas, T., & Dewi, Y. R. (2026). AI Threat or Partner Narratives in Lecturer Student Discourse: A CDA of Polarization and Legitimacy. *Lensa: Kajian Kebahasaan, Kesusastraan, Dan Budaya*, 16(1), 79–102. <https://doi.org/10.26714/lensa.16.1.2026.79-102>

alignment with assessment purposes. (74) (e.g., assignment briefings, classroom discussions, consultation meetings) and written communication (50) (e.g., LMS announcements, rubric explanations, emails, class-group chats). The stance coding resulted in three categories: Threat-oriented statements ($n = 52$), Partner-oriented statements ($n = 35$), and Hybrid/Ambivalent statements ($n = 37$). These frequencies are not used to claim statistical generalization. Rather, they function as an interpretive map for identifying which discourse patterns became more visible and recurrent in the corpus. The higher number of Threat-oriented statements is analytically significant because it shows that GenAI was more frequently discussed in terms of risk, control, academic integrity, and assessment fairness than in terms of learning support.

However, the substantial number of Hybrid/Ambivalent statements indicates that the discourse was not organized as a simple binary opposition. Instead, many statements combined permission with restriction, typically through conditional formulations such as "allowed, but...", "only for...", or "as long as...." This pattern suggests that polarization was stabilized through negotiated boundaries: GenAI could be accepted as a learning partner only when its use was made accountable through disclosure, verification, and alignment with assessment purposes. (74) (e.g., assignment briefings, classroom discussions, consultation meetings) and written communication (50) (e.g., LMS announcements, rubric explanations, emails, class-group chats).

Finding 1: "Threat talk" concentrates in high-stakes assessment moments and is legitimized through authorization and moral evaluation"

A consistent pattern emerges during high-stakes assessment talk (final papers, projects, portfolios): lecturers intensify boundary-making by framing GenAI as a threat to integrity and fairness. The discourse typically relies on authorization (by invoking institutional rules, course policies, rubrics, and sanctions) and on moral evaluation (anchoring the stance in values such as honesty, justice, and respect for peers). Linguistically, threat talk is marked by high modality (e.g., "not allowed," "must," "found out"), categorical prohibitions, and evaluative labels that position AI-mediated writing as morally problematic rather than merely procedurally noncompliant.

L2 (Briefing, spoken): "*Untuk final paper, saya tegaskan tidak boleh pakai ChatGPT untuk menulis paragraf. Kalau ketahuan teksnya hasil AI, itu bukan sekedar salah teknis itu tidak jujur dan merugikan teman yang nulis sendiri.*"
"*For the final paper, I emphasize that students are not allowed to use ChatGPT to write paragraphs. If the text is found to be AI-generated, it is not merely a technical mistake; it is dishonest and unfair to classmates who write their work independently.*" (L2_Briefing_01)

How to Cite (in APA 7th Edition):

erdiawan, R. D., Nurhidayah, Y., Pamungkas, T., & Dewi, Y. R. (2026). AI Threat or Partner Narratives in Lecturer Student Discourse: A CDA of Polarization and Legitimacy. *Lensa: Kajian Kebahasaan, Kesusastraan, Dan Budaya*, 16(1), 79–102. <https://doi.org/10.26714/lensa.16.1.2026.79-102>

In this excerpt, the stance "AI as a Threat" is stabilized by (1) an explicit prohibition ("not allowed") and (2) moral condemnation ("dishonest," "harming friends"). The lecturer's authority is enacted not only through rule-like language but also through the moral positioning of the lecturer as a guardian of fairness—implicitly constructing a legitimate academic identity as one who polices boundaries for the collective good.

Notably, threat talk is also internalized by students, who reproduce the same moral-legal uncertainty, even when describing minimal or supportive uses. In student talk, the key marker is anticipatory risk: students orient to the possibility of being accused, misrecognized, or sanctioned.

S7 (Consultation, spoken): "*Saya takut pakai AI, Bu. Walau cuma buat nyari ide, takutnya nanti dianggap curang, soalnya teman saya pernah dibilang 'itu AI ya' padahal dia cuma paraphrase.*"

"I am afraid to use AI, Ma'am. Even if it is only for finding ideas, I am worried that it will later be considered cheating, because my friend was once accused of using AI even though they only used it for paraphrasing." (S7_Consult_02)

Here, the student does not simply evaluate AI; instead, the student evaluates the 'interpretive environment' an atmosphere where suspicion and ambiguity become part of learning practice. Discursively, "AI as Threat" becomes a social fact not only because lecturers declare it, but because students adapt their behavior to avoid being positioned as unethical writers. This finding highlights that legitimization is interactional: it is co-produced through participants' anticipations of judgment, not merely through official rules.

Finding 2: "Partner talk" becomes prominent in formative learning episodes and is legitimized via rationalization

When interaction shifts from grading and sanctions to learning support (brainstorming, outlining, summarizing readings, language polishing), the discourse often repositions GenAI as a partner. This partner framing is most commonly legitimized through rationalization, the use of practical reasoning that presents GenAI as a tool that increases efficiency, supports comprehension, or scaffolds novice performance. Linguistically, partner talk is marked by conditional permissives ("allowed") and process-oriented verbs ("Help," "Facilitate," "develop"), rather than moral labeling.

L1 (LMS announcement, written): "*ChatGPT boleh digunakan untuk brainstorming topik dan membuat outline. Namun, Anda wajib mengembangkan argumen sendiri, memeriksa akurasi, dan menulis dengan gaya Anda.*"

"ChatGPT may be used for brainstorming topics and creating outlines. However, you are required to develop your own arguments, check the accuracy, and write in your own style." (L1_LMS_05)

How to Cite (in APA 7th Edition):

erdiawan, R. D., Nurhidayah, Y., Pamungkas, T., & Dewi, Y. R. (2026). AI Threat or Partner Narratives in Lecturer Student Discourse: A CDA of Polarization and Legitimacy. *Lensa: Kajian Kebahasaan, Kesusastraan, Dan Budaya*, 16(1), 79–102. <https://doi.org/10.26714/lensa.16.1.2026.79-102>

This statement does two things simultaneously. First, it legitimizes AI as a partner at the idea-generation stage ("may be used"). Second, it protects academic accountability through a set of non-negotiable responsibilities ("mandatory," "check for accuracy," "write in your style"). The lecturer's stance is not "pro-AI" in an unrestricted sense; rather, it is "partner-with-guardrails," suggesting that legitimacy is conditional on student agency and verification.

Students often mirror this rationalized partnership by explicitly narrating verification practices, positioning themselves as responsible users rather than passive recipients of AI output.

S3 (Class discussion, spoken): *"Kalau saya pakai AI untuk meringkas jurnal, itu membantu saya ngerti kerangka artikelnya. Tapi setelah itu saya baca ulang sumbernya, biar nggak ketipu ringkasan yang salah."*

"If I use AI to summarize a journal article, it helps me understand the structure of the article. However, after that, I reread the original source so that I am not misled by an inaccurate summary." (S3_ClassTalk_03)

The student's ("but") functions as a micro-legitimation device: it acknowledges risk while reaffirming responsibility. The student constructs a legitimate identity as an active verifier—someone who uses AI to create cognitive entry points while retaining epistemic control.

Finding 3: Hybrid discourse ("allowed, but...") is the dominant mechanism for stabilizing norms under uncertainty

A major finding is that a large portion of statements operate in a hybrid register: they neither fully reject nor fully endorse GenAI. Instead, they regulate it through conditionality—often realized as concessive constructions ("allowed..., but..., "as long as..., "only for...") and task-based boundaries ("for task," "in which part"). In CDA terms, hybrid discourse functions as an institutional compromise: it reduces conflict by allowing limited use while retaining the authority to sanction perceived overreach.

L4 (Briefing, spoken): *"Saya tidak anti-AI. Boleh pakai untuk cari contoh ungkapan atau cek grammar, tapi kalau AI sudah menyusun paragraf argumentasi, itu sudah melanggar tujuan tugas."* (L4_Briefing_02)

"I'm not anti-AI. It's fine to use it to find examples of expressions or check grammar, but if AI is already composing argumentative paragraphs, that's already violating the purpose of the assignment."

The lecturer here performs identity work ("not anti-AI") to avoid being positioned as technophobic, while simultaneously drawing a line around what "counts" as legitimate academic labor. The boundary is justified via

How to Cite (in APA 7th Edition):

erdiawan, R. D., Nurhidayah, Y., Pamungkas, T., & Dewi, Y. R. (2026). AI Threat or Partner Narratives in Lecturer Student Discourse: A CDA of Polarization and Legitimacy. *Lensa: Kajian Kebahasaan, Kesusastraan, Dan Budaya*, 16(1), 79–102. <https://doi.org/10.26714/lensa.16.1.2026.79-102>

rationalization ("for task"), implying that legitimacy depends on alignment with the pedagogical purpose rather than on the mere presence of AI.

Hybrid regulation is frequently strengthened through written authorization, particularly when lecturers introduce disclosure rules. These statements shift the discussion from "use vs. no use" to "use with transparency," thereby re-centering accountability.

L4 (Email, written): *"Jika Anda menggunakan GenAI, cantumkan bagian mana yang dibantu AI dan jelaskan bagaimana Anda memverifikasi isi. Tanpa disclosure, saya anggap Anda tidak transparan."*

"If you use GenAI, indicate which parts were assisted by AI and explain how you verified the content. Without disclosure, I will consider you not transparent." (L4_Email_01)

Here, legitimation is built through a combined strategy: authorization (a disclosure requirement is framed as a rule) and moral evaluation ("not transparent") as a breach of academic ethics. Importantly, this suggests that "AI as Partner" is institutionally acceptable only when accompanied by traceability, verification, and moralized transparency.

Finding 4: Mythopoesis (cautionary and success stories) powerfully reinforces legitimation and delegitimation

Across the corpus, both lecturers and students frequently use narratives to establish what is acceptable. In van Leeuwen's terms, these are mythopoetic moves: stories that reward certain behaviors and punish others symbolically, making norms memorable and persuasive. The stories often circulate as cautionary warnings about sanctions, lowered grades, or oral defenses, thereby legitimizing stricter control mechanisms. Conversely, students deploy success narratives to normalize AI as a helpful partner while protecting their identity as authentic learners.

L2 (Class talk, spoken): *"Semester lalu ada yang pakai AI buat final. Tulisan terlibat rapi, tapi ketika saya tanya argumennya, dia nggak bisa jelaskan. Akhirnya saya minta oral defense dan nilainya turun drastis."*

"Last semester, someone used AI for the final assignment. The writing looked neat, but when I asked about the argument, they could not explain it. In the end, I asked for an oral defense, and their grade dropped significantly." (L2_ClassTalk_07)

This cautionary story makes two legitimacy claims. First, it defines "real learning" as explainability and ownership of reasoning ("cannot be explained"). Second, it legitimizes surveillance and verification practices ("oral defense") as necessary responses to AI-mediated writing. The story thus transforms a policy stance into lived classroom common sense.

How to Cite (in APA 7th Edition):

erdiawan, R. D., Nurhidayah, Y., Pamungkas, T., & Dewi, Y. R. (2026). AI Threat or Partner Narratives in Lecturer Student Discourse: A CDA of Polarization and Legitimacy. *Lensa: Kajian Kebahasaan, Kesusastraan, Dan Budaya*, 16(1), 79–102. <https://doi.org/10.26714/lensa.16.1.2026.79-102>

Students, however, also produce mythopoetic moves that legitimate AI through constrained agency: AI helps "start," but the student remains the decision-maker.

S11 (Consultation, spoken): "*Saya dulu stuck nulis pembuka. AI bantu saya bikin beberapa opsi hook, terus saya pilih yang paling cocok dan saya revisi besar-besaran. Jadi tetap saya yang putuskan.*"

"I used to get stuck when writing the opening. AI helped me generate several hook options, then I chose the most suitable one and revised it substantially. So, I was still the one who made the final decision." (S11_Consult_01)

The narrative constructs a legitimate partnership by emphasizing selection and revision ("I choose," "major revision," "I still decide"). Discursively, this protects the student's authorial identity while framing GenAI as an assistive brainstorming device.

Finding 5: The central struggle is accountability: "who is responsible?" becomes the core of legitimacy work

Beyond the threat/partner binary, the most persistent tension concerns responsibility: who is accountable for accuracy, originality, and ethical compliance when GenAI is involved? Lecturers repeatedly reassert responsibility as non-transferable ("the writer is responsible, not the tool"), which serves both as authorization (a normative rule of academic culture) and as moral evaluation (blaming "the tool" is portrayed as an illegitimate excuse). Students, meanwhile, foreground inconsistency across lecturers as a source of confusion and anxiety, suggesting that legitimacy is fragile when norms vary between courses.

L5 (Consultation, spoken): "*Kalau AI bilang A, itu bukan berarti benar. Di akademik, yang bertanggung jawab itu penulisnya, bukan alatnya.*"

"If AI says A, it does not necessarily mean that A is correct. In academia, the person who is responsible is the writer, not the tool." (L5_Consult_04)

This statement discursively closes the door on "tool-blame." It re-centers accountability on the student and constructs legitimate academic identity as one who verifies, claims ownership, and bears responsibility for claims made in a text.

S2 (Class group chat, written): "*Jadi kalau pakai AI buat nyusun kalimat itu termasuk curang atau enggak? Aturannya beda-beda tiap dosen. Kita jadi bingung.*" So, *"if we use AI to construct sentences, is that considered cheating or not? The rules are different for each lecturer, so we become confused."* (S2_GroupChat_06)

How to Cite (in APA 7th Edition):

erdiawan, R. D., Nurhidayah, Y., Pamungkas, T., & Dewi, Y. R. (2026). AI Threat or Partner Narratives in Lecturer Student Discourse: A CDA of Polarization and Legitimacy. *Lensa: Kajian Kebahasaan, Kesusastraan, Dan Budaya*, 16(1), 79–102. <https://doi.org/10.26714/lensa.16.1.2026.79-102>

Here, the student frames the problem not as willingness to cheat but as ambiguity of governance across contexts. The "rules differ" claim points to a structural condition that amplifies polarization: inconsistent boundaries make it difficult for students to know which practices are legitimate, and this uncertainty itself becomes a driver of anxiety and self-censorship.

Taken together, the findings show that the discourse of "AI as a Threat vs. Partner" is not merely descriptive; it is performative, producing institutional realities about what students feel safe to do and what lecturers feel obliged to police. "Threat talk" is most dominant when assessment legitimacy is at stake and is often reinforced through authorization, moral evaluation, and cautionary mythopoesis. "Partner talk" becomes more available in formative learning contexts and is legitimized through a rationalization that frames GenAI as scaffolding rather than substitution. Crucially, hybrid discourses pecially the repeated structure "boleh, tapi..." functions as a stabilizing mechanism that allows institutions and classrooms to manage uncertainty by setting conditional boundaries, often through disclosure rules and verification demands. At the center of these discursive moves lies a persistent struggle over accountability and academic identity: whether GenAI use is framed as responsible learning support or illegitimate misrepresentation depends less on the tool itself than on how lecturers and students justify, narrate, and regulate its use in situated interaction.

Discussion

The findings extend current debates on generative AI in higher education by demonstrating that the controversy is less about "using a tool" and more about discursive boundary work, where lecturers and students continually negotiate what counts as legitimate learning, legitimate assistance, and legitimate evidence of authorship in ordinary academic talk. R, Viewed through a CDA lens, these negotiations are not incidental classroom exchanges; they are interactional sites where institutional values such as fairness, responsibility, transparency, and academic integrity are produced and contested. This study contributes to CDA theory by showing that legitimation in GenAI governance is constructed not only through institutional policy discourse but also through everyday lecturer–student interaction. Warnings, permissions, disclosure demands, and conditional approvals function as forms of interactional legitimation that translate abstract institutional values into practical norms of acceptable assistance, authentic authorship, and responsible academic conduct. The study further extends CDA work on legitimation by showing that polarization itself operates as a legitimating mechanism: constructing GenAI as a "threat" legitimizes restriction and control, while constructing it as a "partner" legitimizes scaffolding and conditional acceptance. The recurring pattern of "allowed, but..." is therefore conceptualized as guardrail legitimacy, a form of conditional legitimation through which institutional authority is

How to Cite (in APA 7th Edition):

erdiawan, R. D., Nurhidayah, Y., Pamungkas, T., & Dewi, Y. R. (2026). AI Threat or Partner Narratives in Lecturer Student Discourse: A CDA of Polarization and Legitimacy. *Lensa: Kajian Kebahasaan, Kesusastraan, Dan Budaya*, 16(1), 79–102. <https://doi.org/10.26714/lensa.16.1.2026.79-102>

maintained while limited GenAI use is made acceptable. This explains why GenAI governance often appears unstable or uneven: it is enacted not only through formal policy, but through how policy meanings are translated and negotiated in classroom discourse. These negotiations are not incidental; they are the site where institutional values (fairness, responsibility, transparency) are enacted and contested in real time (Simpson, 2025; van Leeuwen, 2007). This helps explain why institutional responses to GenAI often appear unstable or uneven: governance does not operate only through formal policy, but through how policies are translated into classroom warnings, permissions, and justifications (Dabis & Csáki, 2024; Wang et al., 2024).

Interpreting the results as a legitimacy struggle clarifies the function of polarization. In the dataset, "AI as a threat" and "AI as a partner" did not operate as neutral descriptions of technology but as competing legitimacy frames that authorize some practices while stigmatizing others. Threat-oriented statements tended to protect what can be called the moral economy of assessment fairness, effort, trust in grades by marking GenAI-assisted writing as a potential form of misrepresentation (Cotton et al., 2023; Bretag et al., 2019). This resonates with assessment scholarship, which argues that GenAI destabilizes conventional cues of independent authorship, making it difficult to rely on product-based originality as a defensible indicator of learning (Luo, 2024). Partner-oriented statements, by contrast, framed GenAI as scaffolding and productivity support, aligning with broader syntheses that describe generative conversational AI as simultaneously enabling (feedback, personalization, ideation) and risky (hallucination, bias, opacity, dependency) (Dwivedi et al., 2023; Kasneci et al., 2023). Importantly, the findings show these frames are not simply "opinions"; they are mechanisms for regulating academic conduct and identity.

A key insight is that polarization becomes performative, shaping students' risk-taking and disclosure behavior. When lecturers' discourse centers on moral evaluation (e.g., "cheating," "dishonest"), students tend to anticipate suspicion and recalibrate their practices toward concealment or avoidance, even for uses they believe are pedagogically helpful. This pattern matters because legitimacy becomes fragile when compliance is driven primarily by fear rather than shared understanding; in such conditions, integrity governance may unintentionally encourage "strategic invisibility" rather than responsible transparency (Cotton et al., 2023; Luo, 2024). The multinational student evidence that learners both adopt GenAI for practical academic purposes and simultaneously express anxiety about cheating and regulation provides a macro-level parallel to the micro-level discourse effects observed here (Ravšelj et al., 2025; Yusuf et al., 2024).

The "partner" discourse in the findings is also best understood within the longer arc of digitized academic writing, where students routinely assemble tool ecologies (search engines, grammar checkers, translators, citation managers) to manage complex writing demands. In this light, GenAI is not an isolated rupture but a powerful intensification of existing

How to Cite (in APA 7th Edition):

erdiawan, R. D., Nurhidayah, Y., Pamungkas, T., & Dewi, Y. R. (2026). AI Threat or Partner Narratives in Lecturer Student Discourse: A CDA of Polarization and Legitimacy. *Lensa: Kajian Kebahasaan, Kesusastraan, Dan Budaya*, 16(1), 79–102. <https://doi.org/10.26714/lensa.16.1.2026.79-102>

infrastructures that already mediate what writing "is" in higher education (Zhao et al., 2024). However, the findings suggest a qualitative shift: students describe GenAI as interactive ("dia bisa jelasin lagi," "bisa diajak diskusi"), which strengthens the partner frame beyond older tools and blurs boundaries between assistance, collaboration, and co-production. This helps explain why institutions treat GenAI as legitimacy-sensitive change rather than routine technology adoption, foregrounding transparency, accountability, and human oversight in early responses (Dabis & Csáki, 2024; Wang et al., 2024).

One of the most consequential patterns is the prevalence of hybrid statements permissions bounded by "boleh, tapi..." which the findings interpret as a distinctive form of guardrail legitimacy. Rather than representing a simple compromise, guardrail talk functions as a stabilizing discourse that authorizes limited GenAI use while preserving the assessor's ability to interpret student work as credible evidence of learning. This closely aligns with policy analyses showing an "open but cautious" institutional stance: GenAI may be permitted for certain processes (brainstorming, outlining, language support) if accompanied by disclosure, verification, and clear human responsibility (Dabis & Csáki, 2024; Wang et al., 2024). The guardrails observed in lecturer statements were rarely purely technical ("tool A allowed, tool B banned"); instead, they were framed as ethical and pedagogical criteria (effort, honesty, comprehension, voice), echoing calls to rethink how "originality" is operationalized under machine-assisted writing conditions (Luo, 2024). In language education, this also aligns with emerging arguments that effective governance requires building "ChatGPT literacy" (knowing when, how, and why to use GenAI responsibly) rather than relying solely on prohibition (Ma et al., 2024).

The findings further show that legitimation is accomplished through recognizable discursive strategies consistent with van Leeuwen's framework. Authorization appeared in appeals to policy, rubrics, and institutional rules; moral evaluation mobilized fairness and honesty; rationalization justified GenAI use or restriction via learning goals and practical consequences; and mythopoesis operated through cautionary and exemplary narratives that made norms memorable and enforceable (van Leeuwen, 2007). This pattern is consistent with broader legitimation research showing that legitimacy is produced not only through formal authority but through rhetorical and narrative strategies that naturalize particular interpretations of responsibility and acceptable conduct (Reyes, 2011; Vaara & Tienari, 2008). In this sense, classroom discourse does not simply reflect policy, it performs policy, translating abstract principles into locally compelling stories, labels, and boundary markers (Simpson, 2025).

A central risk highlighted by the findings is the interpretive gap between institutional policy language and classroom-level meaning-making. Even where institutional resources promote responsible use, the local discourse may become highly punitive ("nol toleransi") or highly permissive ("yang penting selesai"), producing uneven norms across courses and

How to Cite (in APA 7th Edition):

erdiawan, R. D., Nurhidayah, Y., Pamungkas, T., & Dewi, Y. R. (2026). AI Threat or Partner Narratives in Lecturer Student Discourse: A CDA of Polarization and Legitimacy. *Lensa: Kajian Kebahasaan, Kesusastraan, Dan Budaya*, 16(1), 79–102. <https://doi.org/10.26714/lensa.16.1.2026.79-102>

increasing student uncertainty about what counts as legitimate assistance (Wang et al., 2024; Simpson, 2025). Teacher-side readiness and prior beliefs also matter: when educators' orientations toward GenAI vary, policy messages can fragment into contradictory classroom regimes, reinforcing student anxiety and strategic behavior (Rahimi & Sevilla-Pavón, 2024). A practical implication is that institutions may need more than "allowed/not allowed" lists; they may require shared communicative templates that help lecturers articulate conditional permissions, request disclosure without stigma, and define "acceptable assistance" with concrete examples that students can interpret consistently (Dabis & Csáki, 2024; Simpson, 2025).

The findings also support a shift in assessment design from policing outputs to evidencing process. Because originality-based assumptions are strained under GenAI, product-only judgments become harder to justify, and overreliance on after-the-fact policing can intensify mistrust and polarization (Luo, 2024). The guardrail discourse observed in the corpus implicitly points toward workable assessment principles: require process evidence (draft trails, decision logs, reflective justifications), embed opportunities for explanation (short oral defenses, annotated revisions), and align tasks with learning outcomes that demand reasoning and disciplinary judgment rather than only fluent textual production (Kasneci et al., 2023; Urquhart & Ngo, 2026; Herdiawan, R. D., Nisa, Z., & Syakuro, P. A., 2025). Such moves do not "solve" GenAI, but they reduce the legitimacy burden placed on textual products alone and better align assessment with visible learning.

Finally, the findings' repeated "detector anxiety" highlights why detection-centered governance can become a legitimacy risk. Reviews of AI content detection tools caution that these systems are not stable foundations for high-stakes decisions due to error rates, shifting model behavior, and context sensitivity (Elkhatat et al., 2023). Equity concerns intensify the legitimacy problem: evidence indicates that GPT detectors can misclassify non-native English writing as AI-generated, raising fairness risks for linguistically diverse students (Liang et al., 2023). In legitimization terms, if enforcement is perceived as unreliable or biased, institutional moral authority is weakened and compliance shifts from ethical commitment to strategic avoidance, widening the interpretive gap and reinforcing the threat/partner polarization.

Overall, this study's contribution lies in relocating the GenAI debate from abstract perceptions and generalized policy claims to the discursive front line of teaching and assessment. By showing how legitimacy is produced through everyday lecturer–student statements via authorization, moral evaluation, rationalization, and narrative the findings clarify why governance succeeds or fails not only through rules but through communicative legitimacy and shared interpretive frames (van Leeuwen, 2007; Wang et al., 2024). In practical terms, sustainable responses require aligning policy, pedagogy, and assessment design with discourse practices that support transparency, protect fairness, and cultivate responsible GenAI literacy

How to Cite (in APA 7th Edition):

herdiawan, R. D., Nurhidayah, Y., Pamungkas, T., & Dewi, Y. R. (2026). AI Threat or Partner Narratives in Lecturer Student Discourse: A CDA of Polarization and Legitimacy. *Lensa: Kajian Kebahasaan, Kesusastraan, Dan Budaya*, 16(1), 79–102. <https://doi.org/10.26714/lensa.16.1.2026.79-102>

without defaulting to stigma or surveillance.

CONCLUSION

This study shows that the debate over GenAI in higher education is not merely about whether lecturers and students use AI, but about how they negotiate legitimacy in academic practice. In high-stakes assessment contexts, lecturers often frame AI as a threat by invoking rules, rubrics, honesty, fairness, and cautionary stories about academic failure. In formative learning contexts, AI is more often framed as a partner because it can support brainstorming, outlining, summarizing, and understanding. However, the dominant pattern is not complete rejection or full acceptance, but conditional permission, expressed through the logic of "allowed, but...". This pattern functions as guardrail legitimacy, in which GenAI use is accepted only when students remain responsible for accuracy, originality, verification, and disclosure. Therefore, universities and lecturers need clearer communication, shared guidance across courses, and assessment practices that focus not only on final products but also on learning processes such as drafts, reflection notes, and brief oral explanations. AI detectors should not be used as the sole basis for judgment, but only as triggers for further review. In this way, GenAI governance can become more transparent, educative, fair, and legitimate.

CRedit AUTHOR STATEMENT

Rama Dwika Herdiawan: Conceptualization, Methodology, Software, Data curation, Writing- Original draft preparation. **Yayah Nurhidayah:** Visualization, Investigation. **Trian Pamungkas:** Supervision, Writing- Reviewing, Validation. **Yaniar Ratna Dewi:** Editing.

ACKNOWLEDGEMENT

The authors would like to express their sincere gratitude to the lecturers and students who volunteered to participate as respondents in this study. Their participation and support were invaluable in providing relevant and in-depth data to support the analysis presented in this research. Special thanks are also extended to the reviewers and proofreaders who provided constructive feedback and suggestions, which significantly contributed to the improvement of this article.

How to Cite (in APA 7th Edition):

herdiawan, R. D., Nurhidayah, Y., Pamungkas, T., & Dewi, Y. R. (2026). AI Threat or Partner Narratives in Lecturer Student Discourse: A CDA of Polarization and Legitimacy. *Lensa: Kajian Kebahasaan, Kesusastraan, Dan Budaya*, 16(1), 79–102. <https://doi.org/10.26714/lensa.16.1.2026.79-102>

REFERENCES

- Braun, V., & Clarke, V. (2006). Using thematic analysis in psychology. *Qualitative Research in Psychology*, 3(2), 77–101. <https://doi.org/10.1191/1478088706qp063oa>
- Bretag, T., Harper, R., Burton, M., Ellis, C., & Rozenberg, P. (2019). Contract cheating: A survey of Australian university staff. *Assessment & Evaluation in Higher Education*, 44(7), 988–1000. <https://doi.org/10.1080/02602938.2019.1640726>
- Cotton, D. R. E., Cotton, P. A., & Shipway, J. R. (2024). Chatting and cheating: Ensuring academic integrity in the era of ChatGPT. *Innovations in Education and Teaching International*, 61(2), 228–239. <https://doi.org/10.1080/14703297.2023.2190148>
- Dabis, A., & Csáki, C. (2024). AI and ethics: Investigating the first policy responses of higher education institutions to the challenge of generative AI. *Humanities and Social Sciences Communications*, 11, Article 1006. <https://doi.org/10.1057/s41599-024-03526-z>
- Dwivedi, Y. K., Kshetri, N., Hughes, L., Slade, E. L., Jeyaraj, A., Kar, A. K., Baabdullah, A. M., Koohang, A., Raghavan, V., Ahuja, M., Albanna, H., Albashrawi, M. A., Al-Busaidi, A. S., Balakrishnan, J., Barlette, Y., Basu, S., Bose, I., Brooks, L., Buhalis, D., ... Wright, R. (2023). So what if ChatGPT wrote it? Multidisciplinary perspectives on opportunities, challenges and implications of generative conversational AI for research, practice and policy. *International Journal of Information Management*, 71, 102642. <https://doi.org/10.1016/j.ijinfomgt.2023.102642>
- Elkhatat, A. M., Elsaid, K., & Almeer, S. (2023). Evaluating the efficacy of AI content detection tools in differentiating between human and AI-generated text. *International Journal for Educational Integrity*, 19, Article 17. <https://doi.org/10.1007/s40979-023-00140-5>
- Herdiawan, R. D., & Windiarti, D. (2025). Teaching Materials for Critical Discourse Analysis And Systematic Functional Linguistics. *Insan Cerdas Bermartabat*, 1-104.

How to Cite (in APA 7th Edition):

erdiawan, R. D., Nurhidayah, Y., Pamungkas, T., & Dewi, Y. R. (2026). AI Threat or Partner Narratives in Lecturer Student Discourse: A CDA of Polarization and Legitimacy. *Lensa: Kajian Kebahasaan, Kesusastraan, Dan Budaya*, 16(1), 79–102. <https://doi.org/10.26714/lensa.16.1.2026.79-102>

- Herdiawan, R. D., Nisa, Z., & Syakuro, P. A. (2025, December). How The Doctoral Students Negotiate Meaning With Ai's Tools? A Case Study On Doctoral Students'writing Pedagogy. In Proceedings of the International Conference and Annual Business Meeting (Vol. 1, No. 1, pp. 69-77).
- Herdiawan, R. D., Fakhrudin, A., Nurhidayat, E., & Nurhadi, K. (2025). Gender stereotypes in visual and verbal texts from government-distributed EFL textbook: Critical discourse analysis. *Journal on English as a Foreign Language*, 15(1), 367-393. <https://doi.org/10.23971/je.fl.v15i1.9357>
- Kasneci, E., Seßler, K., Küchemann, S., Bannert, M., Dementieva, D., Fischer, F., Gasser, U., Groh, G., Günemann, S., Hüllermeier, E., & Kasneci, G. (2023). ChatGPT for good? On opportunities and challenges of large language models for education. *Learning and Individual Differences*, 103, 102274. <https://doi.org/10.1016/j.lindif.2023.102274>
- Liang, W., Yuksekgonul, M., Mao, Y., Wu, E., & Zou, J. (2023). GPT detectors are biased against non-native English writers. *Patterns*, 4(7), 100779. <https://doi.org/10.1016/j.patter.2023.100779>
- Luo, J. (2024). A critical review of GenAI policies in higher education assessment: A call to reconsider the "originality" of students' work. *Assessment & Evaluation in Higher Education*, 49(5), 651–664. <https://doi.org/10.1080/02602938.2024.2309963>
- Ma, Q., Crosthwaite, P., Sun, P. P., & Zou, D. (2024). Developing ChatGPT literacy in language education: A framework and implications. *Computers and Education: Artificial Intelligence*, 7, 100278. <https://doi.org/10.1016/j.cacai.2024.100278>
- Nowell, L. S., Norris, J. M., White, D. E., & Moules, N. J. (2017). Thematic analysis: Striving to meet the trustworthiness criteria. *International Journal of Qualitative Methods*, 16(1), 1–13. <https://doi.org/10.1177/1609406917733847>
- O'Connor, C., & Joffe, H. (2020). Intercoder reliability in qualitative research: Debates and practical guidelines. *International Journal of Qualitative Methods*, 19, 1–13. <https://doi.org/10.1177/1609406919899220>

How to Cite (in APA 7th Edition):

erdiawan, R. D., Nurhidayah, Y., Pamungkas, T., & Dewi, Y. R. (2026). AI Threat or Partner Narratives in Lecturer Student Discourse: A CDA of Polarization and Legitimacy. *Lensa: Kajian Kebahasaan, Kesusastraan, Dan Budaya*, 16(1), 79–102. <https://doi.org/10.26714/lensa.16.1.2026.79-102>

- Rahimi, A., & Sevilla-Pavón, A. (2024). The role of ChatGPT readiness in shaping language teachers' language teaching innovation and meeting accountability: A bisymmetric approach. *Computers and Education: Artificial Intelligence*, 7, 100258. <https://doi.org/10.1016/j.caeai.2024.100258>
- Ravšelj, D., Keržič, D., Tomažević, N., Umek, L., Brezovar, N., Iahad, N. A., Abdulla, A. A., Akopyan, A., Aldana Segura, M. W., AlHumaid, J., Allam, M. F., Alló, M., Andoh, R. P. K., Andronic, O., Arthur, Y. D., Aydın, F., Badran, A., Balbontín-Alvarado, R., Ben Saad, H., ... Aristovnik, A. (2025). Higher education students' perceptions of ChatGPT: A global study of early reactions. *PLOS ONE*, 20(2), e0315011. <https://doi.org/10.1371/journal.pone.0315011>
- Reyes, A. (2011). Strategies of legitimation in political discourse: From words to actions. *Discourse & Society*, 22(6), 781–807. <https://doi.org/10.1177/0957926511419927>
- Simpson, N. H. (2025). Framing AI in higher education: A critical discourse analysis of inclusivity and power in university AI guidance documents. *Educational Linguistics*. Advance online publication. <https://doi.org/10.1515/eduling-2024-0007>
- Urquhart, S., & Ngo, H. (2026). Changing EAP assessment practices in the age of generative artificial intelligence: The case of Scottish higher education institutions. *Journal of English for Academic Purposes*, 69, 102006. <https://doi.org/10.1016/j.jeap.2026.102006>
- Vaara, E., & Tienari, J. (2008). A discursive perspective on legitimation strategies in multinational corporations. *Academy of Management Review*, 33(4), 985–993. <https://doi.org/10.5465/amr.2008.34422019>
- Van Leeuwen, T. (2007). Legitimation in discourse and communication. *Discourse & Communication*, 1(1), 91–112. <https://doi.org/10.1177/1750481307071986>
- Wang, H., Dang, A. N., Wu, Z., & Mac, S. (2024). Generative AI in higher education: Seeing ChatGPT through universities' policies, resources, and guidelines. *Computers and Education: Artificial Intelligence*, 7, 100326. <https://doi.org/10.1016/j.caeai.2024.100326>

How to Cite (in APA 7th Edition):

erdiawan, R. D., Nurhidayah, Y., Pamungkas, T., & Dewi, Y. R. (2026). AI Threat or Partner Narratives in Lecturer Student Discourse: A CDA of Polarization and Legitimacy. *Lensa: Kajian Kebahasaan, Kesusastraan, Dan Budaya*, 16(1), 79–102. <https://doi.org/10.26714/lensa.16.1.2026.79-102>

- Yusuf, B. N., Pervin, N., & Román-González, M. (2024). Generative AI and the future of higher education: A threat to academic integrity or reformation? Evidence from multicultural perspectives. *International Journal of Educational Technology in Higher Education*, 21, Article 21. <https://doi.org/10.1186/s41239-024-00453-6>
- Zhao, X., Cox, A. M., & Cai, L. (2024). ChatGPT and the digitisation of writing. *Humanities and Social Sciences Communications*, 11, Article 482. <https://doi.org/10.1057/s41599-024-02904-x>

How to Cite (in APA 7th Edition):

erdiawan, R. D., Nurhidayah, Y., Pamungkas, T., & Dewi, Y. R. (2026). AI Threat or Partner Narratives in Lecturer Student Discourse: A CDA of Polarization and Legitimacy. *Lensa: Kajian Kebahasaan, Kesusastraan, Dan Budaya*, 16(1), 79–102. <https://doi.org/10.26714/lensa.16.1.2026.79-102>