



Analysis of Semantic and Syntactic Features of Cyberbullying on Instagram and X: Iranian Users

Elahe Kamari ¹ 

Assistant Professor, Department of English Language Teaching, Farhangian University, P.O. Box 14665-889, Tehran, Iran

Abstract

This study investigates the semantic and syntactic patterns of Persian-language cyberbullying among Iranian users on Instagram and X (formerly Twitter). Employing a mixed-methods approach combining qualitative thematic analysis and corpus linguistics, a dataset of n=500 cyberbullying incidents (August–December 2024) was analyzed. The findings reveal that Persian cyberbullying heavily relies on specific semantic targets, with 45% constituting gendered or sexualized slurs, alongside metaphorical dehumanization. Syntactically, perpetrators frequently employ agent-obscuring structures (70%), such as pro-drop and passive constructions, to evade accountability. Furthermore, emphatic punctuation is highly prevalent, with 65% of messages utilizing multiple exclamation marks. The study concludes that cyberbullying in Persian utilizes culturally specific linguistic mechanisms not fully captured by Western-centric algorithms. Developing localized, context-aware artificial intelligence tools is essential for the effective automated detection and mitigation of cyber-aggression in Persian-speaking digital environments.

Keywords: *Content analysis, Cyberbullying, Instagram, Iranian Users, Semantic analysis, Syntactic Structure, X,*

1. ?????? (Corresponding Author)

How to cite: Kamari, E. (2025). Analysis of Semantic and Syntactic Features of Cyberbullying on Instagram and X: Iranian Users. *Language and Linguistics*, 21(42), 203- 232. doi: [10.30465/lsi.2026.51844.1804](https://doi.org/10.30465/lsi.2026.51844.1804)

1. Introduction

The rapid expansion of social media has fundamentally transformed interpersonal communication, simultaneously facilitating a rise in digital hostility. Cyberbullying, generally defined as intentional and repeated harm inflicted through electronic mediums, has emerged as a critical global issue (Smith et al., 2008; Vandebosch & Van Cleemput, 2008). The psychological ramifications of such online aggression are severe, often correlating with elevated anxiety, depression, and even suicidal ideation among victims (Hinduja & Patchin, 2010).

While significant research has been conducted on English-language cyberbullying, non-Western linguistic contexts remain underexplored. In Iran, the widespread adoption of platforms like Instagram and X has generated unique environments for digital interaction. Persian-language cyberbullying exhibits specific cultural and linguistic markers that differentiate it from patterns observed in English contexts. Addressing this gap is vital for understanding the socio-linguistic dimensions of online aggression. Therefore, this study aims to answer the following research questions:

1. What are the dominant semantic targets and thematic categories utilized by Iranian users in cyberbullying incidents?
2. Which specific syntactic and orthographic structures are predominantly employed to convey hostility and minimize perpetrator accountability?
3. How do the distinct architectural features of Instagram and X influence the linguistic manifestation of cyberbullying?

2. Literature Review

Research on cyberbullying has evolved to recognize the complexity of online harassment. Early studies focused on defining the phenomenon and its prevalence (Smith et al., 2008). Subsequent comprehensive meta-analyses, such as that by Kowalski et al. (2014), highlighted the multifaceted nature of cyberbullying, emphasizing the roles of anonymity and platform architecture in facilitating aggressive behavior.

Moving beyond general psychological impacts, researchers have increasingly utilized linguistic and computational approaches to detect online abuse. Dadvar et al. (2013) demonstrated that incorporating user context and specific linguistic features significantly improves the accuracy of machine learning models in identifying cyberbullying. In the Iranian context, Badri et al. (2021) adopted a multidimensional approach, showing

that socio-cultural factors deeply influence the nature of cyber-aggression among Iranian youth. However, detailed syntactic and semantic analyses of Persian cyberbullying texts remain scarce, necessitating targeted corpus-based investigations to decode the specific linguistic strategies employed by Iranian users.

3. METHODOLOGY

This study utilized a descriptive-analytical design, integrating qualitative thematic analysis (Braun & Clarke, 2006) and corpus linguistics (O’Keeffe & McCarthy, 2010). A purposive sample of n=500 text units containing aggressive content was collected (n=250 from public Instagram comments, n=250 from X posts) between August and December 2024. Data collection focused on trending hashtags and controversial public figures. The dataset was annotated and analyzed using MAXQDA 2022 for thematic coding and AntConc for corpus linguistics queries. Reliability was established through dual-coding of a 20% subsample, yielding a high inter-rater agreement (Cronbach’s alpha = 0.89).

4. Results

4.1 Semantic Patterns

The most prevalent pattern was identity-based insult targeting personal characteristics. As shown in Table 1, 45% of offensive vocabulary was gender-based or body-referential, disproportionately targeting women through sexualized slurs. Appearance insults constituted 30%, with higher prevalence on Instagram consistent with its visual orientation. The remaining 25% comprised insults referencing ethnicity (8%), family (7%), and intelligence or mental health (10%).

Table 1
Distribution of Offensive Vocabulary Categories

Category	Frequency (%)
Gender-based / body-referential	45%
Physical appearance	30%
Ethnicity / nationality	8%
Family	7%
Intelligence / mental health / morality	10%

Beyond insult, three further semantic patterns were identified: threats (both disclosure of private information and physical harm); defamation through

false declarative statements presented as objective facts; and dehumanizing metaphor, including animal comparisons (*gāv* (cow), *khuk* (pig), *meymun* (monkey)) and cutting irony deploying honorific titles sarcastically (*Profesor! Fylsuf!*).

4.2 Syntactic and Stylistic Patterns

Thematic analysis revealed five primary semantic targets. The most dominant category was gendered and body-referential insults (45%), followed by attacks on physical appearance (30%). Attacks on intelligence, mental health, or morality constituted 10%, while ethnicity/nationality (8%) and family references (7%) comprised the remainder. Metaphorical dehumanization was a key strategy, frequently utilizing animal comparisons (e.g., “*gāv*” [cow], “*khuk*” [pig]). Sarcasm was also prevalent, often employing honorifics pejoratively (e.g., “*Profesor!*” [Professor!]).

5. Discussion

The findings reveal that cyberbullying among Iranian users on Instagram and X is shaped by both universal patterns of online aggression and features highly specific to the Persian language and cultural context.

The significant prevalence of gender-based and body-referential insults (45%) aligns with the multidimensional nature of cyberbullying globally, where attackers exploit societal vulnerabilities (Kowalski et al., 2014). However, as Badri et al. (2021) observed, the manifestation of these attacks is deeply embedded in the socio-cultural fabric of Iran, where honor, gender roles, and moral reputation carry substantial weight. Metaphorical dehumanization (using animal terms) and sarcastic honorifics serve as calculated strategies to strip targets of their social standing, validating the assertion by Vandebosch and Van Cleemput (2008) that cyberbullying involves complex, intentional power imbalances rather than mere impulsive anger.

Furthermore, the syntactic analysis reveals a strategic linguistic evasion of accountability. By employing agent-obscuring constructions in 70% of the aggressive utterances, perpetrators minimize their visible role in the harassment. The morphological flexibility of the Persian language, particularly its pro-drop nature, facilitates this obfuscation far more readily than grammatically rigid languages. This linguistic complexity supports the findings of Dadvar et al. (2013), who argued that relying solely on profanity lexicons is inadequate; detecting cyberbullying requires a deep understanding of structural and context-specific linguistic features.

Platform differences also played a crucial role, corroborating the findings of Smith et al. (2008) and Kowalski et al. (2014) regarding the impact of digital mediums. Instagram's visually driven architecture fostered sustained appearance-based insults in localized comment threads, whereas X's text-centric, fast-paced environment encouraged rapid, irony-heavy attacks and collective dogpiling.

6. CONCLUSION

This study demonstrates that Persian cyberbullying is not merely direct verbal abuse, but a linguistically calculated behavior characterized by specific semantic targets and syntactic obfuscation. Given the severe psychological distress associated with such victimization (Hinduja & Patchin, 2010), proactive mitigation strategies are essential. The findings emphasize that current, predominantly English-centric detection algorithms are insufficient for identifying Persian cyber-aggression due to their inability to parse nuanced syntactic omissions, sarcastic honorifics, and cultural metaphors. Consequently, integrating Persian corpus linguistics (O'Keeffe & McCarthy, 2010) into the development of localized, context-aware AI tools is imperative for fostering safer digital environments for Iranian users.

Bibliography:

- Badri, R., Kariman, N., Ozgoli, G., & Farsimadan, M. (2021). Sociocultural aspects of cyberbullying among Iranian youth: A qualitative study. *Journal of Interpersonal Violence*, 36(11-12), NP6007-NP6031.
- Braun, V., & Clarke, V. (2006). Using thematic analysis in psychology. *Qualitative Research in Psychology*, 3(2), 77–101.
- Dadvar, M., Trieschnigg, D., Ordelman, R., & de Jong, F. (2013). Improving cyberbullying detection with user context. In *Advances in Information Retrieval: 35th European Conference on IR Research* (pp. 693–696). Springer.
- Hinduja, S., & Patchin, J. W. (2010). Bullying, cyberbullying, and suicide. *Archives of Suicide Research*, 14(3), 206–221.
- Kowalski, R. M., Giumetti, G. W., Schroeder, A. N., & Lattanner, M. R. (2014). Bullying in the digital age: A critical review and meta-analysis of cyberbullying research among youth. *Psychological Bulletin*, 140(4), 1073–1137.
- O'Keeffe, A., & McCarthy, M. (Eds.). (2010). *The Routledge handbook of corpus linguistics*. Routledge.

- Smith, P. K., Mahdavi, J., Carvalho, M., Fisher, S., Russell, S., & Tippett, N. (2008). Cyberbullying: Its nature and impact in secondary school pupils. *Journal of Child Psychology and Psychiatry*, 49(4), 376–385.
- Vandebosch, H., & Van Cleemput, K. (2008). Defining cyberbullying: A Qualitative research into the perceptions of youngsters. *CyberPsychology & Behavior*, 11(4), 499–503.



بررسی ویژگی‌های معنایی و نحوی زورگویی سایبری در اینستاگرام و ایکس: کاربران ایرانی

الهه کمری 

استادیار، گروه آموزش زبان انگلیسی، دانشگاه فرهنگیان، صندوق پستی ۸۸۹-۱۴۶۶۵، تهران، ایران

چکیده

زورگویی سایبری به عنوان یکی از چالش‌های آسیب‌زای فضای مجازی، با گسترش روزافزون استفاده از شبکه‌های اجتماعی در سطح جهانی و به ویژه در ایران، اهمیت مضاعفی یافته است. این پدیده می‌تواند آثار مخربی بر قربانیان داشته باشد. پژوهش حاضر با هدف تحلیل ویژگی‌های زبانی، معنایی (مانند واژگان توهین‌آمیز، استعاره‌های تحقیرآمیز، مضامین غالب) و نحوی (مانند ساختار جملات، کاربرد علائم نگارشی، حذف یا ابهام عامل) در محتوای زورگیرانه تولید شده توسط کاربران ایرانی در دو سکوی محبوب اینستاگرام و ایکس (تویتر سابق) انجام پذیرفت. داده‌های پژوهش شامل مجموعه‌ای از ۵۰۰ پست و نظرات مرتبط با آن‌ها (۲۵۰ نمونه از هر سکو) بود که با استفاده از روش نمونه‌گیری هدفمند و بر اساس معیارهای مشخص نشان‌دهنده پتانسیل زورگویی سایبری، طی بازه زمانی مشخصی گردآوری شدند. تحلیل داده‌ها با بهره‌گیری از روش تحلیل محتوای کیفی با رویکرد رمزگذاری مضمون و با پشتیبانی نرم‌افزار تحلیل داده‌های کیفی MAXQDA صورت گرفت. همچنین، از ابزارها و مفاهیم زبان‌شناسی پیکره‌ای برای شناسایی الگوهای تکرارشونده استفاده شد. یافته‌های کلیدی نشان داد که واژگان با بار معنایی جنسیت‌محور (حدود ۴۵ درصد از کل واژگان توهین‌آمیز شناسایی شده) و همچنین استفاده مکرر و تأکیدی از علائم نگارشی خاص (به‌ویژه علامت تعجب چندگانه !!! در ۶۵ درصد موارد) نقش برجسته‌ای در ساخت و تشدید فضای پرخاشگرانه دارند. علاوه بر این، مشاهده شد که در بخش قابل توجهی از محتوا (حدود ۷۰ درصد)، ساختارهای نحوی فاقد فعل معلوم یا با فاعل مبهم (مانند جملات مجهول یا حذف فاعل) به‌منظور پنهان‌سازی هویت فرد زورگو یا کاهش مسئولیت مستقیم به کار گرفته می‌شود. این نتایج بر ضرورت حیاتی توسعه و بومی‌سازی الگوریتم‌ها و ابزارهای تشخیص خودکار زورگویی سایبری تأکید دارد؛ ابزارهایی که قادر باشند نشانگرهای زبانی خاص فرهنگ و زبان فارسی در جامعه ایرانی را به دقت شناسایی و تحلیل کنند.

کلیدواژه: زورگویی سایبری، تحلیل معنایی، ساختار نحوی، کاربران ایرانی، اینستاگرام، ایکس، تحلیل محتوا.

۱- مقدمه

عصر حاضر را بی‌شک می‌توان عصر دیجیتال و غلبه فضای مجازی بر شئون مختلف زندگی فردی و اجتماعی دانست. گسترش اینترنت و به تبع آن، ظهور و توسعه شبکه‌های اجتماعی^۱ انقلابی در الگوهای ارتباطی، تعاملات اجتماعی، دسترسی به اطلاعات و حتی شکل‌گیری هویت افراد ایجاد کرده است (کاپلان^۲ و هانلین^۳، ۲۰۱۰). ایران نیز از این قاعده مستثنی نبوده و شاهد رشد فزاینده کاربران فضای مجازی، به ویژه در میان نسل جوان، بوده است. سکوهایی نظیر اینستاگرام با ماهیت بصری‌محور و ایکس (توییتر سابق) با ویژگی بلاگ‌نویسی کوچک و سرعت بالای انتشار اطلاعات، به بخش جدایی‌ناپذیری از زندگی روزمره میلیون‌ها ایرانی تبدیل شده‌اند (بدری و همکاران، ۲۰۲۱). این سکوها فرصت‌های بی‌شماری برای ارتباط، یادگیری، سرگرمی و کنشگری اجتماعی فراهم آورده‌اند؛ اما در عین حال، بستری برای ظهور و گسترش پدیده‌های آسیب‌زا و چالش‌های جدید نیز شده‌اند.

یکی از مهم‌ترین و نگران‌کننده‌ترین این چالش‌ها، پدیده «زورگویی سایبری»^۴ است. زورگویی سایبری به طور کلی به استفاده عمدی و مکرر از فناوری‌های اطلاعاتی و ارتباطی مانند اینترنت، تلفن همراه، شبکه‌های اجتماعی برای آسیب رساندن، تحقیر، تهدید یا ارباب دیگران اطلاق می‌شود (اسمیت^۵ و همکاران، ۲۰۰۸؛ هیندوجا^۶ و بومان^۷، ۲۰۱۳). این پدیده، برخلاف زورگویی سنتی که محدود به مکان و زمان خاصی بود، می‌تواند در هر زمانی از شبانه‌روز، مخاطبان گسترده‌تری داشته باشد (به دلیل قابلیت انتشار سریع محتوا)، و به دلیل امکان پنهان ماندن هویت فرد زورگو (ناشناس بودن)، از پیچیدگی و آثار مخرب بیشتری برخوردار باشد (اسلونجی^۸ و اسمیت، ۲۰۰۸). قربانیان زورگویی سایبری ممکن است دچار اضطراب، افسردگی، افت تحصیلی، انزوای اجتماعی، و در موارد حاد، حتی افکار خودکشی شوند (کووالسکی^۹ و همکاران، ۲۰۱۴). با توجه به محبوبیت بالای اینستاگرام و ایکس در میان کاربران ایرانی، این دو سکو به عرصه‌های اصلی وقوع زورگویی سایبری در این بافت فرهنگی تبدیل شده‌اند و شناخت ابعاد مختلف این پدیده در آن‌ها ضرورت می‌یابد.

-
1. Social Networking Sites-SNSs
 2. Kaplan, A. M.
 3. Haenlein, M.
 4. cyberbullying
 5. Smith, P. K.
 6. Hinduja, S.
 7. Bauman, S.
 8. Slonje, R.
 9. Kowalski, R. M.

پدیده زورگویی سایبری در سطح بین‌المللی موضوع مطالعات گسترده‌ای در رشته‌های مختلف از جمله روانشناسی، جامعه‌شناسی، علوم ارتباطات و علوم رایانه بوده است (توکوناگا^۱ و بومان، ۲۰۱۳؛ و پاچین^۲ و هیندوجا، ۲۰۱۲). بسیاری از این پژوهش‌ها بر شناسایی عوامل خطر، ویژگی‌های شخصیتی زورگوها و قربانیان، پیامدهای روانشناختی و اجتماعی، و راهکارهای پیشگیری و مداخله متمرکز بوده‌اند. همچنین، در حوزه علوم رایانه، تلاش‌های زیادی برای توسعه الگوریتم‌های یادگیری ماشین جهت تشخیص خودکار محتوای زورگیرانه صورت گرفته است (دادور و همکاران، ۲۰۱۳؛ سالوو^۳ و همکاران، ۲۰۲۰). با این حال، بخش عمده‌ای از تحقیقات موجود، به ویژه در زمینه تحلیل محتوای زبانی و تشخیص خودکار، بر روی داده‌های مربوط به جوامع انگلیسی‌زبان انجام شده است. این در حالی است که زبان، به عنوان ابزار اصلی ارتکاب زورگویی سایبری متنی، عمیقاً تحت تأثیر بافت فرهنگی، اجتماعی و هنجارهای زبانی خاص هر جامعه قرار دارد. شیوه‌های بیان پرخاشگری، انواع توهین‌ها، استفاده از کنایه‌ها، استعاره‌ها، و حتی ساختارهای نحوی که برای انتقال معنای تهدیدآمیز یا تحقیرآمیز به کار می‌روند، می‌توانند از زبانی به زبان دیگر و از فرهنگی به فرهنگ دیگر تفاوت‌های چشمگیری داشته باشند (لونیگرو^۴ و همکاران، ۲۰۱۵).

در ایران، اگرچه مطالعاتی در زمینه شیوع زورگویی سایبری و برخی جنبه‌های جامعه‌شناختی و روانشناختی آن انجام شده است (مانند محمدی و همکاران، ۱۴۰۱؛ فلاحی و همکاران، ۱۴۰۱)، اما شکاف پژوهشی قابل توجهی در زمینه بررسی دقیق و نظام‌مند نشانگرهای زبانی خاص این پدیده در زبان فارسی و در بستر شبکه‌های اجتماعی پرکاربرد ایرانیان مشهود است. به عبارت دیگر، کمتر پژوهشی به این موضوع پرداخته است که کاربران ایرانی چگونه از ظرفیت‌های معنایی (واژگان، اصطلاحات، استعاره‌ها) و نحوی (ساختار جمله، علائم نگارشی) زبان فارسی برای اعمال زورگویی سایبری در سکوهایی مانند اینستاگرام و ایکس بهره می‌برند. عدم شناخت این ویژگی‌های زبانی بومی، می‌تواند منجر به ناکارآمدی ابزارهای تشخیصی توسعه‌یافته بر اساس داده‌های غیرفارسی و همچنین دشواری در طراحی مداخلات آموزشی و فرهنگی مؤثر شود.

هدف اصلی این پژوهش، واکاوی و توصیف نظام‌مند ویژگی‌های معنایی و نحوی محتوای متنی زورگیرانه تولید شده توسط کاربران ایرانی در سکوهایی اینستاگرام و ایکس است. این پژوهش در پی آن است که با تحلیل داده‌های واقعی از این دو شبکه اجتماعی، الگوهای زبانی

-
1. Tokunaga, R. S.
 2. Patchin, J. W.
 3. Salaw, S.
 4. Lonigro, A.

پراکارد در ارتکاب زورگویی سایبری به زبان فارسی را شناسایی کرده و در نهایت، چارچوبی مقدماتی برای درک بهتر و شناسایی دقیق‌تر این پدیده در بافت فرهنگی-زبانی خاص ایران ارائه دهد. به طور مشخص‌تر، این مطالعه به دنبال پاسخگویی به پرسش‌های زیر است:

۱. پراکاردترین الگوهای معنایی (شامل واژگان توهین‌آمیز، مضامین، استعاره‌ها و کنایه‌ها) در محتوای زورگیرانه کاربران ایرانی در اینستاگرام و ایکس کدامند؟ آیا می‌توان دسته‌بندی مشخصی از این الگوها ارائه داد؟ آیا تفاوتی در الگوهای معنایی بین دو سکو وجود دارد؟
۲. ساختارهای نحوی (مانند طول جملات، نوع جملات، استفاده از وجه معلوم و مجهول، کاربرد علائم نگارشی خاص) چگونه به ساخت، تشدید و انتقال پیام‌های پرخاشگرانه و زورگیرانه در این محتواها کمک می‌کنند؟ آیا الگوهای نحوی مشخصی برای پنهان‌سازی هویت یا افزایش تأثیرگذاری کلام زورگو به کار می‌رود؟

۲. پیشینه پژوهش

مرور ادبیات پژوهشی مرتبط با زورگویی سایبری نشان می‌دهد که این حوزه از منظر رشته‌ها و رویکردهای مختلفی مورد بررسی قرار گرفته است. برای ارائه تصویری جامع و در عین حال متمرکز بر اهداف این پژوهش، پیشینه مطالعاتی را می‌توان در سه گروه اصلی دسته‌بندی کرد: پژوهش‌های متمرکز بر جنبه‌های روانشناختی و اجتماعی، تحقیقات مرتبط با تحلیل محتوای متنی (عمدتاً در زبان انگلیسی)، و مطالعات انجام‌شده در بافت ایران.

بخش قابل توجهی از تحقیقات اولیه و جاری در حوزه زورگویی سایبری به بررسی ویژگی‌های فردی و اجتماعی درگیر در این پدیده اختصاص یافته است. پژوهشگران تلاش کرده‌اند تا عوامل خطر و محافظت‌کننده مرتبط با زورگویی و قربانی شدن در فضای مجازی را شناسایی کنند. مطالعاتی مانند پژوهش اسمیت و همکاران (۲۰۰۸ و ۲۰۱۰) به بررسی شیوع، انواع و تأثیرات زورگویی سایبری در میان نوجوانان در کشورهای مختلف پرداخته‌اند. آن‌ها دریافته‌اند که زورگویی سایبری اغلب با زورگویی سنتی همپوشانی دارد، اما ویژگی‌های منحصر به فرد فضای مجازی (مانند گستردگی مخاطب و ناشناسی) می‌تواند آن را تشدید کند. تحقیقات دیگر بر ویژگی‌های شخصیتی و روانشناختی افراد درگیر متمرکز شده‌اند. برای مثال، رونیونز^۱ و همکاران (۲۰۱۳) و بریور^۲ و کرسلیک^۳ (۲۰۱۵) نشان داده‌اند که سطوح پایین همدلی، پرخاشگری بالا، و مشکلات رفتاری با احتمال بیشتر ارتکاب زورگویی سایبری مرتبط است. از سوی دیگر، قربانیان زورگویی سایبری بیشتر در معرض خطر ابتلا به مشکلات سلامت

1. Runions, K.

2. Brewer, G.

3. Kerslake, J.

روان مانند افسردگی، اضطراب، عزت نفس پایین و افکار خودکشی قرار دارند (هیندوجا و پاچین، ۲۰۱۰؛ کووالسکی و همکاران، ۲۰۱۴). عوامل اجتماعی مانند نظارت والدین، جو مدرسه و حمایت همسالان نیز به عنوان متغیرهای تأثیرگذار بر شیوع و پیامدهای زورگویی سایبری شناسایی شده‌اند (فستل^۱ و همکاران، ۲۰۱۳). این دسته از مطالعات، اگرچه به طور مستقیم به تحلیل زبانی نمی‌پردازند، اما درک عمیق‌تری از زمینه روانی و اجتماعی که زورگویی سایبری در آن رخ می‌دهد، فراهم می‌کنند و به فهم انگیزه‌ها و تأثیرات احتمالی نهفته در پس انتخاب‌های زبانی کمک می‌کنند.

تحقیقات قابل توجهی بر تحلیل محتوای متنی زورگویی سایبری و توسعه سیستم‌های تشخیص خودکار آن متمرکز شده‌اند. این مطالعات بیشتر از پیکره‌های داده^۲ جمع‌آوری شده از شبکه‌های اجتماعی مانند توییتر، فیسبوک، یوتیوب و اینستاگرام (در بافت انگلیسی‌زبان) استفاده کرده‌اند. پژوهشگرانی مانند پچین (پچین، ۲۰۱۵؛ هیندوجا و پچین، ۲۰۱۵) به تحلیل محتوای کیفی پیام‌های زورگیرانه پرداخته و مضامین رایج مانند توهین، شایعه‌پراکنی، تهدید، طرد اجتماعی و جعل شناسایی کرده‌اند. مطالعات دیگر با استفاده از رویکردهای کمی و محاسباتی، به استخراج ویژگی‌های زبانی نشانگر زورگویی سایبری پرداخته‌اند. این ویژگی‌ها شامل موارد زیر بوده‌اند:

ویژگی‌های واژگانی: حضور کلمات توهین‌آمیز^۳، کلمات با بار عاطفی منفی، ضمائر دوم شخص مانند «تو» برای خطاب مستقیم به قربانی (دادور^۴ و همکاران، ۲۰۱۳؛ نهرو^۵ و همکاران، ۲۰۱۲).

ویژگی‌های نحوی و سبکی: طول جملات، استفاده از حروف بزرگ^۶، تکرار حروف یا علائم نگارشی (مانند !!! یا ???) برای تأکید یا ابراز هیجان شدید (رینولدز^۷ و همکاران، ۲۰۱۱؛ دیناکار و همکاران^۸، ۲۰۱۲).

1. Festl, R.
2. corpora
3. profanity
4. Dadvar, M.
5. Nahar, V.
6. capitalization
7. Reynolds, K.
8. Dinakar, K.

ویژگی‌های معنایی و کاربردشناختی: شناسایی کنش‌های گفتاری خاص مانند تهدید یا توهین، استفاده از نشانگرهای سمی بودن^۱ در متن (وان هی^۲ و همکاران، ۲۰۱۵؛ واسیم^۳ و هووی^۴، ۲۰۱۶).

این تحقیقات پایه‌های مهمی برای درک چگونگی بازنمایی زبانی زورگویی سایبری فراهم کرده‌اند و منجر به توسعه ابزارهای اولیه برای تشخیص آن شده‌اند. با این حال، همانطور که پیشتر اشاره شد، اتکای بیشتر آن‌ها بر داده‌های انگلیسی‌زبان، قابلیت تعمیم‌پذیری مستقیم یافته‌ها و ابزارهای حاصل به زبان‌ها و فرهنگ‌های دیگر، از جمله فارسی، را با محدودیت مواجه می‌سازد. تفاوت در ساختار زبان، هنجارهای فرهنگی ابراز پرخاشگری، و حتی شیوه‌های استفاده از سکوها می‌تواند منجر به بروز الگوهای زبانی متفاوتی شود.

مطالعات محدود پیرامون زورگویی سایبری در ایران در سال‌های اخیر، توجه پژوهشگران ایرانی نیز به پدیده زورگویی سایبری جلب شده است. با این حال، بیشتر مطالعات انجام‌شده در این زمینه، بر جنبه‌های جامعه‌شناختی، روانشناختی و شیوع‌شناسی متمرکز بوده‌اند. به عنوان مثال، پژوهش محمدی و همکاران (۱۴۰۱) به بررسی رابطه بین استفاده از شبکه‌های اجتماعی و تجربه زورگویی سایبری در میان دانشجویان پرداخت و عوامل اجتماعی مؤثر بر آن را بررسی کرد. مطالعه فلاحی و همکاران (۱۴۰۱) شیوع و انواع زورگویی سایبری را در بین دانش‌آموزان شهر اردبیل سنجید و پیامدهای روانشناختی آن را گزارش داد. تحقیقات دیگری نیز به بررسی نقش متغیرهایی مانند جنسیت، سن، مدت زمان استفاده از اینترنت، و نظارت والدین بر درگیری در زورگویی سایبری پرداخته‌اند (زارع و همکاران، ۱۳۹۹).

این مطالعات نقش مهمی در ترسیم وضعیت زورگویی سایبری در ایران و برجسته کردن اهمیت آن به عنوان یک مسئله اجتماعی ایفا کرده‌اند. با این وجود، تحلیل زبانی محتوای زورگیرانه در این مطالعات مغفول مانده است. به عبارت دیگر، پژوهش‌ها کمتر به این موضوع پرداخته‌اند که پرخاشگری و قلدری در فضای مجازی فارسی‌زبان دقیقاً با چه واژگان، اصطلاحات، ساختارهای جمله‌ای و شگردهای زبانی‌ای اعمال می‌شود. مرور پیشینه نشان می‌دهد که علیرغم حجم قابل توجه پژوهش‌ها در سطح بین‌المللی و آغاز توجه به این پدیده در ایران، نبود یک پژوهش زبان‌محور نظام‌مند که به طور خاص بر شناسایی و تحلیل الگوهای معنایی و نحوی زورگویی سایبری در میان کاربران فارسی‌زبان در سکوها پر کاربرد اینستاگرام و ایکس تمرکز کند، به وضوح احساس می‌شود. مطالعه حاضر برای پر کردن همین خلأ پژوهشی

1. toxicity

2. Van Hee, C.

3. Waseem, Z.

4. Hovy, D.

است. این پژوهش با اتخاذ رویکردی کیفی و بهره‌گیری از ابزارهای تحلیل محتوا و مفاهیم زبان‌شناسی پیکره‌ای، تلاش دارد تا ویژگی‌های زبانی منحصربه‌فرد زورگویی سایبری در بافت فرهنگی-زبانی ایران را آشکار سازد. این امر نه تنها موجب تقویت دانش نظری افراد را در این حوزه می‌شود، بلکه می‌تواند به طراحی راهکارهای عملی مؤثرتر برای مقابله با این پدیده در جامعه ایران کمک کند.

۳. مبانی نظری

برای تحلیل عمیق ویژگی‌های معنایی و نحوی زورگیری سایبری در زبان فارسی و درک چگونگی کارکرد زبان در این پدیده، پژوهش حاضر از دو چارچوب نظری مکمل بهره می‌گیرد: نظریه کنش گفتاری^۱ و زبان‌شناسی پیکره‌ای^۲. نظریه کنش گفتاری، که بنیان آن توسط جان آستین^۳ (۱۹۶۲) گذاشته شد و سپس توسط جان سرل^۴ (۱۹۶۹، ۱۹۷۶) بسط یافت، بر این ایده محوری استوار است که زبان صرفاً برای توصیف جهان به کار نمی‌رود، بلکه افراد از طریق گفتار (یا نوشتار)، کنش‌هایی را انجام می‌دهند. آستین سه سطح از کنش را در هر گفته^۵ تشخیص داد: (۱) کنش بیانی یا لفظی^۶ که عبارت است از بیان جمله با تبعیت از قوانین دستور زبان برای انتقال معنا؛ (۲) کنش غیربیانی^۷ یا منظوری که با نیت گوینده مرتبط است؛ و (۳) کنش تأثیری^۸ که عواقبی دارد که بر شنونده تأثیر می‌گذارد.

سرل (۱۹۷۶) در نظریه‌اش، به‌خصوص بر روی کنش غیربیانی تمرکز کرده و آن را به پنج نوع مختلف تقسیم‌بندی کرده است: (۱) کنش‌های اظهاری^۹ که به کنش‌های گفتاری گفته می‌شوند که گوینده را به حقیقت گزاره بیان‌شده متعهد می‌کند. این عمل به گوینده امکان می‌دهد تا باورهای خود را بیان کند، که می‌تواند توسط شنونده به‌صورت درست یا نادرست ارزیابی شود. گزارش‌دهی^{۱۰}، تأکید^{۱۱} و نتیجه‌گیری^{۱۲} نمونه‌هایی از اظهاری‌ها هستند. (۲) کنش‌های ترغیبی^{۱۳} توسط گوینده برای تحمیل یک اقدام بر شنونده هستند، مانند

-
1. Speech Act Theory
 2. Corpus Linguistics
 3. Austin, J. L.
 4. Searle, J.
 5. utterance
 6. locutionary act
 7. illocutionary act
 8. perlocutionary act
 9. representative
 10. reporting
 11. asserting
 12. concluding
 13. directive

درخواست^۱ و سؤال^۲. ۳) کنش‌های تعهدی^۳ به گوینده امکان می‌دهند که خود را متعهد به انجام کاری کند. این عمل در وعده‌ها^۴، ضمانت‌ها^۵ و سوگندها^۶ استفاده می‌شود. ۴) کنش‌های عاطفی^۷ بر حالت‌های روان‌شناختی تأثیر می‌گذارند و در تبریک‌ها^۸، عذرخواهی^۹، شکایات^{۱۰}، احساس رضایت^{۱۱} و نارضایتی‌ها^{۱۲} به کار می‌روند. به عبارت دیگر، این نوع کنش گفتاری احساسات گوینده را بیان می‌کند. ۵) کنش‌های اعلامی^{۱۳} تغییراتی در وضعیت نهادی ایجاد می‌کنند. این عمل به‌طور عمده در موقعیت‌های رسمی مانند صدور حکم^{۱۴}، استعفا^{۱۵} و اعلام جنگ^{۱۶} استفاده می‌شود.

نظریه کنش گفتاری چارچوب مفهومی قدرتمندی برای تحلیل زورگیری سایبری فراهم می‌کند، زیرا زورگیری اساساً از طریق انجام کنش‌های گفتاری آسیب‌زا صورت می‌گیرد. محتوای زورگیرانه صرفاً مجموعه‌ای از کلمات نیست، بلکه شامل کنش‌های ضمنی گفتاری مانند تهدید کردن (یک نوع تعهدی یا ترغیبی)، توهین کردن (یک نوع اظهاری)، تحقیر کردن (اظهاری)، اتهام زدن (اخباری با نیت آسیب‌زا)، شایعه‌پراکنی (اخباری با نیت آسیب‌زا)، و طرد کردن (که می‌تواند از طریق ترکیبی از کنش‌ها انجام شود) است. این نظریه کمک می‌کند تا از سطح توصیف صرف واژگان و جملات فراتر رفته و به نقش^{۱۷} و قصد^{۱۸} زبانی در متن زورگیرانه توجه شود که چگونه انتخاب‌های واژگانی و ساختارهای نحوی خاص برای تحقق بخشیدن به این کنش‌های گفتاری پرخاشگرانه به کار گرفته می‌شوند. تأثیرات احتمالی (کنش‌های تأثیری) این گفته‌ها بر قربانی را بهتر درک کنیم (مانند ایجاد ترس، اضطراب، یا احساس حقارت). بنابراین، در تحلیل داده‌ها، شناسایی کنش‌های گفتاری غالب در محتوای زورگیرانه و بررسی ابزارهای زبانی تحقق آن‌ها در کانون توجه پژوهش حاضر خواهد بود.

-
1. requesting
 2. questioning
 3. commissive
 4. promise
 5. guaranties
 6. vow
 7. expressive
 8. congratulation
 9. excuse
 10. complaint
 11. like
 12. dislike
 13. declarative
 14. sentencing
 15. resignation
 16. declaring war
 17. function
 18. intent

زبان‌شناسی پیکره‌ای یک رویکرد روش‌شناختی به مطالعهٔ زبان است که بر تحلیل داده‌های زبانی واقعی و گسترده که «پیکره»^۱ نامیده می‌شود با استفاده از ابزارهای رایانه‌ای تأکید دارد (مکانری^۲ و ویلسون^۳، ۲۰۰۱؛ بیکر^۴، ۲۰۰۶). این رویکرد به جای اتکا بر شهود زبانی پژوهشگر یا مثال‌های ساختگی، بر الگوهای واقعی کاربرد زبان در بافت طبیعی تمرکز می‌کند. برخی از اصول و روش‌های کلیدی زبان‌شناسی پیکره‌ای عبارتند از:

- **استفاده از پیکره‌های بزرگ و معتبر:** جمع‌آوری حجم قابل توجهی از داده‌های زبانی مرتبط با موضوع پژوهش.

- **تحلیل کمی و کیفی:** استفاده از روش‌های آماری برای شناسایی الگوهای فراوانی (مانند واژگان پرتکرار، همایندهای کلیدی) و سپس تحلیل کیفی این الگوها در بافت.

- **ابزارهای تحلیل:** استفاده از نرم‌افزارهای پیکره‌ای برای جستجو، استخراج فهرست واژگان^۵، محاسبهٔ فراوانی، شناسایی کلمات کلیدی^۶، تحلیل همایندها^۷، و بررسی خطوط همخوانی^۸ که نمونه‌های کاربرد یک واژه یا عبارت را در بافت نشان می‌دهد.

رویکرد زبان‌شناسی پیکره‌ای برای مطالعهٔ حاضر بسیار مناسب است زیرا این امکان را فراهم می‌کند تا الگوهای معنایی (مانند واژگان توهین‌آمیز پرتکرار، مضامین غالب) و نحوی (مانند ساختارهای جمله‌ای رایج، کاربرد خاص علائم نگارشی) که به طور نظام‌مند در حجم زیادی از داده‌های زورگیری سایبری فارسی تکرار می‌شوند، به صورت تجربی و عینی شناسایی شوند. با استفاده از تحلیل فراوانی و کلمات کلیدی، می‌توان واژگان و مفاهیمی را که به طور خاص در گفتمان زورگیری سایبری برجسته هستند (در مقایسه با یک پیکره مرجع از زبان فارسی عمومی در دسترس، یا حداقل در مقایسه با بخش‌های غیرزورگیرانهٔ داده‌ها) را مشخص نمود. تحلیل همایندها می‌تواند نشان دهد که چه کلماتی معمولاً با هم به کار می‌روند و چگونه این ترکیب‌ها به ساخت معنای توهین‌آمیز یا تهدیدآمیز کمک می‌کنند. بررسی خطوط همخوانی این امکان را فراهم می‌کند تا کاربرد واژگان و ساختارهای کلیدی را در بافت دقیق‌تری بررسی کرده و تحلیل کیفی عمیق‌تری ارائه شود. این رویکرد به شناسایی ویژگی‌هایی مانند استفاده

1. corpus
2. McEnery, T.
3. Wilso, A.
4. Baker, P.
5. wordlist
6. keyword
7. collocation
8. concordance line

تأکیدی از علائم نگارشی یا الگوهای ساختاری خاص (مانند جملات بدون فعل معلوم) که در نگاه اول ممکن است پنهان بمانند، کمک می‌کند.

در این پژوهش، نظریه کنش گفتاری و زبان‌شناسی پیکره‌ای به صورت مکمل به کار گرفته می‌شوند. زبان‌شناسی پیکره‌ای ابزارها و روش‌های لازم برای شناسایی الگوهای زبانی پرتکرار (معنایی و نحوی) در داده‌های واقعی زورگیری سایبری را فراهم می‌کند. سپس، نظریه کنش کمک می‌کند تا نقش و کارکرد این الگوهای شناسایی شده را در انجام کنش‌های گفتاری پرخاشگرانه و زورگیرانه تفسیر شوند. به عبارت دیگر، از روش‌های پیکره‌ای برای یافتن اینکه «چه چیزی» در زبان زورگیری سایبری فارسی برجسته است، استفاده می‌شود و از نظریه کنش گفتاری برای فهمیدن «چگونه» و «چرا» این عناصر زبانی برای آسیب رساندن به دیگران به کار می‌روند. این تلفیق، امکان تحلیلی جامع و چندبعدی از پدیده مورد مطالعه را فراهم می‌آورد.

۴. روش

پژوهش حاضر از نظر هدف، یک پژوهش اکتشافی-توصیفی است، زیرا به دنبال شناسایی و توصیف ویژگی‌های زبانی یک پدیده نسبتاً کمتر مطالعه‌شده در بافت زبان فارسی است. از نظر رویکرد گردآوری و تحلیل داده‌ها، این پژوهش مبتنی بر رویکرد کیفی است که با تحلیل‌های کمی توصیفی تکمیل می‌شود. طرح اصلی پژوهش، تحلیل محتوای کیفی با رویکرد تحلیل مضمون^۱ است. تحلیل مضمون روشی انعطاف‌پذیر برای شناسایی، تحلیل و گزارش الگوها یا مضامین در داده‌های کیفی است (براون^۲ و کلارک^۳، ۲۰۰۶). این روش به پژوهشگر امکان می‌دهد تا داده‌ها را به صورت نظام‌مند کدگذاری کرده و مضامین اصلی مرتبط با پرسش‌های پژوهش را استخراج و تفسیر کند. در این مطالعه، تمرکز بر شناسایی مضامین مرتبط با ویژگی‌های معنایی (مانند انواع توهین، مضامین پرخاشگرانه) و نحوی (مانند الگوهای ساختاری، استفاده از نشانگرهای سبکی) در محتوای زورگیرانه خواهد بود. همچنین، از اصول و ابزارهای زبان‌شناسی پیکره‌ای به صورت کمکی برای شناسایی الگوهای پرتکرار واژگانی و ساختاری استفاده می‌شود.

جامعه آماری این پژوهش شامل کلیه محتوای متنی (پست‌ها و نظرات عمومی) تولید شده توسط کاربران ایرانی در دو سکوی شبکه اجتماعی اینستاگرام و ایکس (توییتر سابق) در بازه زمانی فروردین تا اردیبهشت ۱۴۰۳ است که پتانسیل طبقه‌بندی به عنوان زورگیری سایبری

1. Thematic Analysis
2. Braun
3. Clarke

را دارند. با توجه به گستردگی و عدم امکان دسترسی به کل این جامعه، تمرکز بر روی محتوای تولید شده توسط کاربران قرار گرفت که بر اساس اطلاعات پروفایل یا محتوای تولیدی، در گروه سنی تقریبی ۱۸ تا ۳۵ سال قرار می‌گرفتند. این بازه سنی به دلیل نرخ بالای استفاده از این دو سکو و همچنین گزارش‌های پیشین مبنی بر شیوع بیشتر زورگیری سایبری در میان جوانان و بزرگسالان جوان انتخاب شد (پاتچین و هیندوجا، ۲۰۱۲). تأکید بر «کاربران ایرانی» به معنای کاربرانی است که به زبان فارسی محتوا تولید می‌کنند و هویت ایرانی خود را (به‌طور مستقیم یا غیرمستقیم) در پروفایل یا محتوایشان بروز می‌دهند.

با توجه به ماهیت اکتشافی پژوهش و دشواری شناسایی قطعی موارد زورگیری سایبری پیش از تحلیل، از روش نمونه‌گیری هدفمند^۱ استفاده شد. در این روش، پژوهشگران بر اساس قضاوت و دانش خود، نمونه‌هایی را انتخاب می‌کنند که بیشترین اطلاعات را در مورد پدیده مورد مطالعه فراهم کنند (پاتون^۲، ۲۰۱۵). معیارهای اصلی برای انتخاب پست‌ها و نظرات جهت ورود به نمونه عبارت بودند از: عمومی بودن محتوا (عدم نیاز به دنبال کردن یا مجوز برای مشاهده). نگارش محتوا به زبان فارسی. وجود نشانگرهای اولیه احتمالی زورگیری سایبری، مانند: استفاده از زبان تهاجمی، توهین‌آمیز یا تهدیدآمیز. هدف قرار دادن یک فرد یا گروه خاص به صورت منفی. وقوع در بستر یک مجادله یا بحث برخاسته. دریافت پاسخ‌ها یا واکنش‌هایی که نشان‌دهنده برداشت منفی یا آزاردهنده بودن محتوا توسط دیگر کاربران باشد. حجم نمونه نهایی شامل ۵۰۰ واحد تحلیل (پست یا نظر) بود که به طور مساوی بین دو سکو تقسیم شد: ۲۵۰ نمونه از اینستاگرام و ۲۵۰ نمونه از ایکس. انتخاب این تعداد با هدف دستیابی به تنوع کافی در داده‌ها و امکان شناسایی الگوهای تکرارشونده و رسیدن به اشباع نسبی داده‌ها^۳ در چارچوب محدودیت‌های زمانی و منابع پژوهش صورت گرفت. هر واحد تحلیل می‌توانست یک پست اصلی یا یک نظر باشد که به عنوان محتوای زورگیرانه بالقوه شناسایی شد. همچنین، نظرات یا پاسخ‌های مرتبط با آن واحد تحلیل نیز برای درک بهتر بافت، مورد بررسی قرار گرفتند.

گردآوری داده‌ها در دو مرحله و با استفاده از ابزارهای زیر انجام شد:

نرم‌افزارهای استخراج داده از وب^۴: برای جمع‌آوری اولیه حجم زیادی از محتوای عمومی از صفحات و هشتک‌های پربازدید و مستعد بحث‌وجدل در اینستاگرام و ایکس (با رعایت کامل ملاحظات اخلاقی و تمرکز صرف بر داده‌های عمومی)، از ابزارهای استاندارد وب اسکرپینگ

1. purposive sampling
2. Patton, M. Q.
3. data saturation
4. web scraping tool

(مانند کتابخانه‌های پایتون نظیر *requests* و *BeautifulSoup* استفاده شد. این ابزارها به جمع‌آوری متون پست‌ها، نظرات، و اطلاعات فرامتنی عمومی (مانند تعداد می‌پسندم یا توییت دوباره، بدون ذخیره اطلاعات هویتی کاربران) کمک کردند. جستجوها با استفاده از کلیدواژه‌های فارسی مرتبط با مجادلات رایج، موضوعات حساسیت‌برانگیز اجتماعی، و همچنین رصد صفحات و حساب‌های کاربری شناخته‌شده برای تولید محتوای بحث‌برانگیز انجام شد. بازه زمانی جمع‌آوری داده‌ها از مرداد ماه ۱۴۰۳ تا دی ماه ۱۴۰۳ حدود ۶ ماه در نظر گرفته شد تا تغییرات احتمالی در روندهای زبانی نیز تا حدی پوشش داده شود.

معیارهای غربالگری و پرسشنامه کدگذاری محقق‌ساخته: پس از جمع‌آوری اولیه، داده‌ها توسط تیم پژوهش (شامل حداقل دو کدگذار آموزش‌دیده) بر اساس معیارهای نمونه‌گیری هدفمند به دقت غربالگری شدند تا ۵۰۰ نمونه نهایی انتخاب شوند. برای اطمینان از پایایی فرآیند انتخاب و کدگذاری اولیه، یک فرم کدگذاری توسط پژوهشگر طراحی شد. این فرم شامل چک‌لیستی از نشانگرهای اولیه زورگیری سایبری (مانند وجود توهین مستقیم، تهدید، تمسخر و...) و فضایی برای ثبت اولیه مشاهدات بود. پایایی بین کدگذاران^۱ برای مرحله غربالگری و همچنین برای کدگذاری مضامین نهایی محاسبه شد. آلفای کرونباخ = ۰/۸۹ نشان‌دهنده پایایی بسیار خوب و همسویی بالای کدگذاران در ارزیابی و دسته‌بندی اولیه محتوا است.

تحلیل داده‌های متنی گردآوری‌شده طی مراحل زیر و با استفاده از نرم‌افزار تحلیل داده‌های کیفی *MAXQDA* (نسخه ۲۰۲۲) انجام شد:

آشنایی با داده‌ها: مطالعه مکرر متن نمونه‌ها برای دستیابی به درک عمیق از محتوا و بافت. **کدگذاری باز:**^۲ خوانش خط‌به‌خط یا بخش‌به‌بخش داده‌ها و اختصاص کدهای اولیه به قسمت‌هایی از متن که ویژگی‌های معنایی یا نحوی مرتبط با زورگیری سایبری را نشان می‌دادند. این کدها تا حد امکان نزدیک به متن اصلی و توصیفی بودند (مانند: «توهین جنسیتی»، «استفاده از تعجب مکرر»، «جمله بدون فاعل مشخص»، «تهدید غیرمستقیم»).

ایجاد دسته‌ها و مضامین:^۳ کدهای اولیه مشابه و مرتبط با هم، در دسته‌های^۴ وسیع‌تری گروه‌بندی شدند. سپس، با بررسی روابط بین دسته‌ها و بازخوانی مداوم داده‌ها و کدها، مضامین اصلی مرتبط با پرسش‌های پژوهش (الگوهای معنایی و الگوهای نحوی) استخراج شدند. این

1. inter-coder reliability
2. open Coding
3. axial coding & thematic analysis
4. category

فرآیند به صورت رفت و برگشتی و با اصلاح مداوم کدها و مضامین انجام شد تا بهترین تناسب با داده‌ها حاصل شود (براون و کلارک، ۲۰۰۶).

تحلیل آماری توصیفی: پس از نهایی شدن کدگذاری مضامین، از قابلیت‌های نرم‌افزار *MAXQDA* برای محاسبه فراوانی و درصد رخداد کدهای مختلف، دسته‌ها و مضامین استفاده شد. این تحلیل‌های کمی توصیفی به ارائه تصویری کلی از شیوع الگوهای شناسایی شده کمک کرد (مثلاً درصد واژگان توهین‌آمیز جنسیت‌محور یا درصد استفاده از ساختارهای نحوی خاص). نتایج این بخش در قالب جداول و نمودارها (مانند موارد ذکر شده در بخش یافته‌ها) ارائه می‌شود.

تحلیل پیکره‌ای کمکی: برای شناسایی واژگان پرتکرار خاص گفتمان زورگیری و همایندهای کلیدی، از ابزارهای پایه زبان‌شناسی پیکره‌ای (مانند ایجاد فهرست فراوانی واژگان، جستجوی همایندها) در نرم‌افزار *MAXQDA* یا نرم‌افزارهای پیکره‌ای دیگر مانند *AntConc* بر روی پیکره کوچک‌شده ۵۰۰ نمونه‌ای استفاده شد تا یافته‌های کیفی با شواهد کمی بیشتری پشتیبانی شوند.

در تمام مراحل انجام این پژوهش، اصول اخلاقی کار با داده‌های انسانی و محتوای برخاسته از پست‌ها و نظراتی استفاده شد که به صورت عمومی در دسترس بودند و نیازی به کسب اجازه برای مشاهده آن‌ها نبود. همچنین، کلیه اطلاعات شناسایی‌کننده کاربران (مانند نام کاربری، نام واقعی، عکس پروفایل) پیش از تحلیل، از داده‌ها حذف یا جایگزین با کدهای شناسایی شد تا هویت افراد محفوظ بماند. داده‌های جمع‌آوری شده به صورت امن نگهداری شده و صرفاً برای اهداف این پژوهش مورد استفاده قرار گرفتند. با توجه به ماهیت حساس و بالقوه آزاردهنده محتوای زورگیرانه، پژوهشگران با احتیاط و رعایت اصول حرفه‌ای با داده‌ها مواجه شدند و از نقل قول مستقیم موارد بسیار شدید یا قابل ردیابی در گزارش نهایی خودداری کرده و در صورت لزوم، مثال‌ها را اندکی تغییر دادند تا ضمن حفظ معنا، قابلیت شناسایی از بین برود.

۵. یافته‌ها

تحلیل محتوای کیفی ۵۰۰ نمونه گردآوری شده از پست‌ها و نظرات کاربران ایرانی در اینستاگرام و ایکس، الگوهای مشخصی را در استفاده از زبان برای اعمال زورگیری سایبری آشکار ساخت. یافته‌های اصلی در دو بخش الگوهای معنایی و الگوهای نحوی ارائه می‌شوند. تحلیل‌های آماری توصیفی نیز برای نشان دادن فراوانی این الگوها به کار گرفته شده‌اند.

۵-۱. الگوهای معنایی غالب در محتوای زورگیرانه

تحلیل معنایی بر شناسایی واژگان کلیدی، مضامین تکرارشونده، و شگردهای معنایی مانند استعاره و کنایه متمرکز بود. کدگذاری مضمونی منجر به شناسایی چندین تم اصلی در این بخش شد:

۵-۱. مضمون ۱: توهین و تحقیر مبتنی بر هویت و ویژگی‌های فردی

این شایع‌ترین الگوی معنایی بود. توهین‌ها اغلب هویت جنسیتی، گرایش جنسی، قومیت، ظاهر فیزیکی، وضعیت خانوادگی، یا توانایی‌های ذهنی فرد را هدف قرار می‌دادند. واژگان توهین‌آمیز جنسیت‌محور: این دسته از واژگان بیشترین فراوانی را داشتند. جدول ۱ توزیع انواع اصلی واژگان توهین‌آمیز را نشان می‌دهد. همانطور که مشاهده می‌شود، ۴۵ درصد از کل واژگان توهین‌آمیز شناسایی شده، مستقیماً به جنسیت یا بدن افراد (به ویژه زنان) ارجاع داشتند و با هدف تحقیر جنسی یا زیر سوال بردن ارزش فرد بر اساس جنسیتش به کار می‌رفتند. مثال‌هایی از این دسته شامل استفاده از الفاظ رکیک جنسی، نسبت دادن صفات منفی کلیشه‌ای به یک جنسیت، و تحقیر بر اساس نقش‌های جنسیتی سنتی بود. توهین‌های مرتبط با ظاهر: دومین دسته پرکاربرد (حدود ۳۰ درصد) شامل توهین‌هایی بود که ظاهر فیزیکی فرد (وزن، چهره، پوشش) را مسخره یا تحقیر می‌کردند. این نوع توهین در سکوی دیداری محور اینستاگرام شیوع بیشتری داشت. سایر توهین‌ها: دسته‌های دیگر شامل توهین‌های قومیتی، توهین به خانواده، اتهامات اخلاقی بی‌اساس، و برچسب‌های مرتبط با هوش یا سلامت روان (مانند «احمق»، «دیوانه») بودند که مجموعاً حدود ۲۵ درصد باقی‌مانده را تشکیل می‌دادند.

جدول ۱: توزیع انواع اصلی واژگان توهین‌آمیز در نمونه تحلیل‌شده

مثال‌های مضمون	درصد فراوانی تقریبی	دسته واژگان توهین‌آمیز
الفاظ رکیک: کاربرد واژگان صریح جنسی (مانند «ک...کش»، «ت...م سگ»); تحقیر جنسیتی: برچسب‌زنی با واژگانی نظیر «هرزه»، «فاحشه»، «لاشی» و «پلنگ».	۴۵٪	جنسیت‌محور و ارجاع به بدن

ادامه جدول ۱:

مثال‌های مضمون	درصد فراوانی تقریبی	دسته واژگان توهین‌آمیز
وزن: تمسخر با الفاظ «بشکه»، «خوک»؛ چهره: حمله به ظاهر با واژگان «دماغ عملی»، «پلاستیکی»؛ پوشش و نقص: استفاده از صفات «خز»، «دهاتی تیپ»، «کچل» و تمسخر نقص عضو.	۳۰٪	مرتبط با ظاهر فیزیکی
بازتولید کلیشه‌های منفی قومیتی؛ استفاده تحقیرآمیز از واژه «افغانی» یا «عرب...»؛ استفاده از واژه «دهاتی» به معنای بی‌فرهنگ؛ تمسخر لهجه.	۸٪	مرتبط با قومیت/ملیت
دشنام مستقیم به والدین (مانند «پدرسگ»); زیر سؤال بردن اصالت و مشروعیت تولد با واژگانی نظیر «حرام‌زاده»، «بی‌بوته» و «لقمه حرام»	۷٪	مرتبط با خانواده
جنون: برچسب‌زنی با کلمات «روانی»، «شیزوفرنی»، «تیمارستانی»، «روانی»؛ نادانی: تحقیر ضریب هوشی با الفاظ «عقب‌مانده»، «اسکل»، «بی‌سواد»؛ بی‌اخلاقی: استفاده از صفات «وطن‌فروش» و «بی‌شرف».	۱۰٪	مرتبط با هوش/سلامت روان/اخلاق

۵-۱-۲. مضمون ۲: تهدید و ارباب

اگرچه کمتر از توهین مستقیم شایع بود، اما موارد قابل توجهی از تهدید (مستقیم یا غیرمستقیم) نیز مشاهده شد. این تهدیدها می‌توانست شامل تهدید به افشای اطلاعات خصوصی^۱ («آدرست رو دارم...»، «عکس‌ها رو پخش می‌کنم» یا «دایرکت رو چک کن تا بفهمی»)، تهدید به آسیب فیزیکی (اغلب به صورت اغراق‌آمیز یا غیرمستقیم) (می‌دم خفتت

1. doxing

کنن»، «پیدات کنم زنده نمی‌دارم» یا «توی خیابون حواست باشه» بسامد بالایی داشتند)، یا تهدید به آسیب رساندن به آبرو یا موقعیت اجتماعی فرد («کاری می‌کنم پیجت بسته شه» یا «آبروت رو توی فامیل می‌برم) باشد.

۵-۱-۳. مضمون ۳: شایعه‌پراکنی و افترا

انتشار اطلاعات نادرست یا تحریف‌شده با هدف تخریب وجهه قربانی، الگوی معنایی دیگری بود که بیشتر حول محور مسائل اخلاقی و مالی می‌چرخید. شواهد متنی نشان می‌دهد که مهاجمان با جملاتی خبری و قطعی سعی در واقعی جلوه دادن دروغ‌ها داشتند؛ عباراتی مانند «این پول‌ها رو از راه ... درآورده»، «خودم دیدم که با فلانی رابطه داره»، «مدرک تحصیلیش جعلیه» یا «این آدم کلاهبرداره». این افتراها اغلب در بخش نظرات فرسته‌ها به صورت هماهنگ تکرار می‌شد تا ذهنیت مخاطبان عام را تغییر دهد.

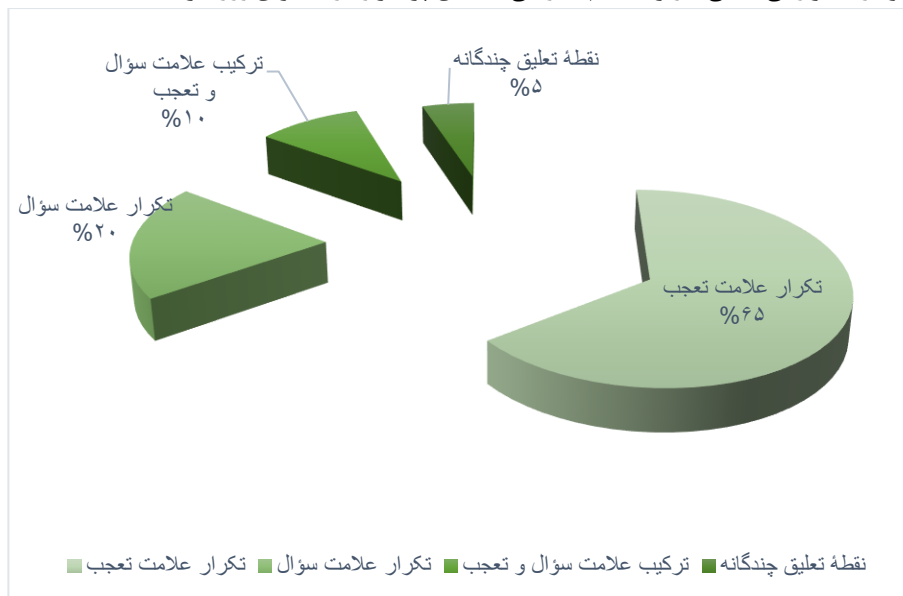
۵-۱-۴. مضمون ۴: استفاده از استعاره‌ها و کنایه‌های تحقیرآمیز

زورگیران سایبری علاوه بر دشنام صریح، از زبان مجازی برای فرار از مسئولیت و افزایش بار تحقیر استفاده می‌کردند: استعاره‌های حیوانی: تشبیه مستقیم قربانی به حیوانات برای انسانیت‌زدایی؛ مانند خطاب کردن فرد با واژگانی چون «گاو» (برای نادانی)، «میمون» (برای ظاهر)، «خوک» (برای پرخوری/چاقی) و «طوطی» (برای تقلید کورکورانه). کنایه‌های گزنده و القاب معکوس: استفاده از القاب محترمانه در بافت تمسخرآمیز؛ نظیر کاربرد واژه «پروفسور!» یا «فیلسوف!» برای کسی که اشتباهی مرتکب شده، یا عبارت «باشه تو خوبی» و «خوشگل خانوم» (برای مردان) با هدف تحقیر جنسیتی و شخصیتی.

۵-۲. الگوهای نحوی و سبکی در خدمت پرخاشگری

یکی از برجسته‌ترین یافته‌های نحوی، استفاده نامتعارف از علائم نگارشی برای بازنمایی فریاد زدن یا خشم شدید بود: علامت تعجب چندگانه (!!!!): با فراوانی ۶۵٪، بیشترین کاربرد را داشت. جملاتی مانند «خفه شو!!!!!!» یا «چقدر احمقی تو!!!!» نشان می‌داد که تکرار علامت تعجب برای افزایش شدت توهین و انتقال بار هیجانی خشم ضروری انگاشته می‌شود. علامت سوال چندگانه (؟؟؟): در ۲۰٪ موارد مشاهده شد که کارکرد آن پرسش واقعی نبود، بلکه برای ناباوری تحقیرآمیز یا به چالش کشیدن قربانی به کار می‌رفت (برای مثال: «واقعاً مدرک داری؟؟؟؟» یا «چی زدی؟؟؟؟»). ترکیب علائم و کشیدگی حروف: استفاده از ترکیب «!؟!» برای ایجاد سردرگمی و فشار روانی، و یا کشیدن حروف برای تمسخر لهجه یا تأکید (مانند: «بُـــــــرو بابا») از دیگر شگردهای سبکی پرتکرار در نمونه‌ها بود.

نمودار ۱. فراوانی نسبی کاربرد علائم نگارشی تأکیدی پرتکرار در محتوای زورگیرانه



۵-۲-۱. ساختارهای نحوی فاقد عامل مشخص^۱

یافته مهم دیگر، شیوع بالای ساختارهایی بود که هویت فاعل یا عامل کنش منفی را پنهان یا مبهم می‌کردند. حدود ۷۰ درصد از جملات حاوی توهین یا اتهام در نمونه تحلیل‌شده، از ساختارهای فاقد فعل معلوم با فاعل مشخص استفاده می‌کردند. این امر از طریق راهکارهای نحوی زیر حاصل می‌شد:

الف) استفاده از جملات مجهول و غیرشخصی: به جای بیان صریح فاعل (من)، گوینده با حذف عاملیت خود، توهین را به یک واقعیت عینی یا ضرورت خارجی تبدیل می‌کرد. استفاده از جملاتی نظیر «باید ریپورت شی» (به جای «من ریپورت می‌کنم»)، «دیده شده که...»، یا «اینجا جای تو نیست، باید حذف بشی». این ساختار، حمله را نه به عنوان خصومت شخصی، بلکه به عنوان یک مجازات جمعی یا قانون نانوشته جلوه می‌دهد.

ب) حذف قرینه فاعل: با توجه به ویژگی ضمیرانداز بودن زبان فارسی، در بسیاری موارد جملات توهین‌آمیز بدون ذکر فاعل بیان می‌شدند تا تمرکز به طور کامل بر صفت تحقیرآمیز باشد. جملاتی که به طور مستقیم با گزاره شروع می‌شوند؛ مانند «[تو] خیلی احمقی»، «[تو]

¹ agentless construction

مایه ننگی»، «حیف نون» یا «اصلاً شعور نداری». حذف «تو» در اینجا باعث می‌شود صفت نسبت داده شده (احمق، ننگ) برجسته‌تر از مخاطب به نظر برسد.

ج) استفاده از ساختارهای اسمی یا عبارت‌های بدون فعل: به جای جملات کامل، استفاده از «برچسب‌های کلامی رواج داشت. این عبارات مانند ضربات سریع و کوتاه عمل می‌کردند. استفاده از عباراتی مانند «یه عقده‌ای بدبخت»، «دلچک مجازی»، «وطن فروش خائن» یا «ندید بدید» بدون هیچ فعل یا مقدمه‌ای، تنها با هدف تخریب سریع شخصیت.

د) ارجاع به فاعل جمع یا نامشخص: استفاده از افعال جمع یا ارجاع به «دیگران» برای کاهش مسئولیت فردی و ایجاد فشار اجتماعی. جملاتی که ادعای اجماع دارند؛ مانند «همه می‌دونن که کلاهبرداری»، «کل ایران ازت متنفره»، «میگن خیلی لاشی هستی» یا «ملت دیگه بیدار شدن، حنای تو رنگی نداره». در اینجا فرد زورگو حمله خود را پشت نقاب «نظر عمومی» یا «حقیقت پذیرفته‌شده» پنهان می‌کند.

جملات کوتاه، بریده‌بریده و تلگرافی: به‌ویژه در سکوی ایکس (تویتر) و فرسته‌های پرخاشگرانه اینستاگرام، زورگیری از طریق جملات کوتاه و کوبنده صورت می‌گرفت تا حس تهاجم را تشدید کند. جملات دستوری تک‌کلمه‌ای یا کوتاه نظیر «خفه»، «بمیر بابا»، «گمشو»، «ساکت شو» یا «فقط سکوت». این ایجاز کلامی، فرصت دفاع را از قربانی می‌گیرد و حکم قطعی صادر می‌کند.

استفاده از حروف بزرگ در واژگان انگلیسی: اگرچه زبان فارسی فاقد حروف بزرگ و کوچک است، اما در بافت دوزبانه فضای مجازی ایران، استفاده از کلمات انگلیسی با حروف بزرگ برای «فریاد زدن مجازی» یا تأکید بر انگ‌زنی مشاهده شد. نوشتن واژگانی مانند *(FAKE)* «فیک/تقلب»، *(LOOSER)* «بازنده» در «تو یه آدم *FAKE* هستی» یا فقط *BLOCK*.

تفاوت‌های احتمالی بین سکوها

اگرچه الگوهای کلی معنایی و نحوی در هر دو سکو مشابهت‌های زیادی داشتند، اما تفاوت‌های ظریفی نیز مشاهده شد. در اینستاگرام به دلیل ماهیت بصری، توهین‌های مرتبط با ظاهر در این سکو شایع‌تر بود. همچنین، به دلیل فضای فرسته‌گذاری طولانی‌تر، گاهی زورگیری به شکل بحث‌های طولانی و کش‌دار با تبادل توهین‌های مکرر دیده می‌شد. ایکس: به دلیل محدودیت کاراکتر و سرعت بالای انتشار، زورگیری اغلب به شکل حملات کوتاه، گزنده، و با استفاده بیشتر از کنایه، هشتگ‌های توهین‌آمیز، و منشن کردن مستقیم افراد

برای بسیج حمله^۱ صورت می‌گرفت. تهدید به افشای اطلاعات یا تحرکات هماهنگ‌شده در ایکس بیشتر مشاهده شد. این یافته‌ها نشان می‌دهند که زورگیری سایبری در میان کاربران ایرانی در اینستاگرام و ایکس، از الگوهای زبانی نسبتاً مشخصی پیروی می‌کند که ترکیبی از انتخاب‌های واژگانی هدفمند (به‌ویژه جنسیت‌محور) و بهره‌گیری از شگردهای نحوی و سبکی (مانند تأکید با علائم نگارشی و پنهان‌سازی عامل) است.

۶. بحث و نتیجه‌گیری

این پژوهش با هدف واکاوی ویژگی‌های معنایی و نحوی محتوای زورگیرانه تولید شده توسط کاربران ایرانی در سکوها اینستاگرام و ایکس انجام شد. یافته‌های به‌دست‌آمده از تحلیل کیفی و کمی ۵۰۰ نمونه محتوا، الگوهای زبانی مشخصی را آشکار ساخت که در این بخش مورد بحث و تفسیر قرار گرفته و با پیشینه پژوهش و چارچوب نظری مرتبط می‌شوند. همچنین، دلالت‌های کاربردی، محدودیت‌ها و پیشنهادهایی برای پژوهش‌های آتی ارائه خواهد شد.

الگوهای معنایی (پاسخ به پرسش ۱): یافته‌ها نشان داد که پرکاربردترین الگوی معنایی در محتوای زورگیرانه کاربران ایرانی، توهین و تحقیر مبتنی بر هویت و ویژگی‌های فردی است. برجستگی خاص واژگان توهین‌آمیز جنسیت‌محور (۴۵ درصد) و توهین‌های مرتبط با ظاهر (۳۰ درصد)، نشان‌دهنده تأثیر عمیق هنجارها و کلیشه‌های فرهنگی مرتبط با جنسیت و ظاهر در شکل‌دهی به زبان پرخاشگری برخط در جامعه ایران است. این یافته با برخی مطالعات بین‌المللی که به نقش جنسیت در زورگیری سایبری پرداخته‌اند همخوانی دارد (اولوس^۲، ۲۰۱۲)، اما شدت و تمرکز خاص بر این نوع توهین‌ها ممکن است بازتابی از حساسیت‌ها و مسائل فرهنگی ویژه بافت ایران باشد. سایر الگوهای معنایی مانند تهدید، شایعه‌پراکنی و استفاده از زبان غیرمستقیم (کنایه و استعاره) نیز، اگرچه با فراوانی کمتر، نقش مهمی در زرادخانه زبانی زورگیران سایبری ایفا می‌کنند. این تنوع در کنش‌های گفتاری پرخاشگرانه (توهین، تهدید، افترا) با چارچوب نظری کنش گفتاری (آستین و سرل) همخوانی دارد و نشان می‌دهد که زورگیران از طیف وسیعی از ابزارهای زبانی برای اعمال کنش‌های آسیب‌زا استفاده می‌کنند.

الگوهای نحوی و سبکی (پاسخ به پرسش ۲): تحلیل نحوی نشان داد که ساختار جملات و انتخاب‌های سبکی نیز به طور هدفمند برای انتقال و تشدید پرخاشگری به کار می‌روند. استفاده بسیار زیاد از علائم نگارشی تأکیدی (!!! و ؟؟؟)، که در ۶۵ درصد و ۲۰ درصد موارد مشاهده شد،

1. mobbing

2. Olweus, D.

راهکاری سبکی برای جبران فقدان نشانه‌های فرازبانی (مانند لحن صدا یا حالات چهره) در ارتباط متنی و انتقال شدت هیجان منفی (خشم، تمسخر) است. این یافته با برخی مطالعات پیشین در زبان انگلیسی (رینولدز و همکاران، ۲۰۱۱) همخوانی دارد، اما فراوانی بالای آن در نمونه فارسی قابل توجه است. یافته کلیدی دیگر، یعنی شیوع بالای ساختارهای نحوی فاقد عامل مشخص (۷۰ درصد) از طریق مجهول‌سازی، حذف فاعل یا استفاده از جملات اسمی، شگردی زبانی برای کاهش مسئولیت‌پذیری و پنهان‌سازی هویت زورگو است. این الگو، که شاید در زبان فارسی به دلیل انعطاف‌پذیری ساختاری آن راحت‌تر قابل استفاده باشد، لایه دیگری از پیچیدگی را به تشخیص زورگیری سایبری می‌افزاید، زیرا پرخاشگری به شکلی غیرمستقیم‌تر و مبهم‌تر بیان می‌شود. این الگوها نشان می‌دهند که تحلیل زورگیری سایبری نباید صرفاً بر روی واژگان کلیدی متمرکز باشد، بلکه باید به ساختارهای نحوی و نشانگرهای سبکی نیز توجه ویژه‌ای داشت. رویکرد زبان‌شناسی پیکره‌ای در شناسایی این الگوهای پرتکرار نحوی و سبکی بسیار مؤثر بود.

این پژوهش ضمن تأیید برخی یافته‌های کلی مطالعات بین‌المللی (مانند اهمیت توهین و تهدید)، به دلیل تمرکز بر زبان فارسی و بافت فرهنگی ایران، یافته‌های جدید و متمایزی را ارائه می‌دهد. در حالی که مطالعات پیشین در ایران عمدتاً بر جنبه‌های جامعه‌شناختی و روانشناختی متمرکز بودند (مانند محمدی و همکاران، ۱۴۰۱)، این پژوهش نخستین گام‌ها را در تحلیل نظام‌مند زبانی زورگیری سایبری فارسی در سکوه‌های محبوب برداشت. شناسایی برجستگی توهین‌های جنسیت‌محور و مرتبط با ظاهر، و همچنین شیوع بالای ساختارهای نحوی بدون عامل مشخص، یافته‌هایی هستند که به طور خاص به درک این پدیده در بافت ایران کمک می‌کنند و شکاف پژوهشی شناسایی‌شده در زمینه تحلیل زبان‌محور را تا حدی پر می‌کنند. این یافته‌ها همچنین اهمیت فراتر رفتن از تحلیل‌های مبتنی بر کلیدواژه‌های ساده (که در بسیاری از ابزارهای تشخیص خودکار اولیه به کار می‌روند) و لزوم توجه به الگوهای معنایی و نحوی پیچیده‌تر و وابسته به بافت را برجسته می‌سازد.

این مطالعه به غنای دانش در حوزه زبان‌شناسی اجتماعی، تحلیل گفتمان انتقادی، و مطالعات ارتباطات میان‌فرهنگی کمک می‌کند. نشان می‌دهد که چگونه زبان به عنوان ابزاری برای اعمال قدرت و خشونت در فضای مجازی به کار گرفته می‌شود و چگونه هنجارهای فرهنگی و ساختارهای زبانی در این فرآیند نقش دارند. همچنین، لزوم توجه به ابعاد زبانی خاص هر فرهنگ در مطالعه پدیده‌های برخط جهانی مانند زورگیری سایبری را مورد تأکید قرار می‌دهد.

یافته‌های این پژوهش به‌ویژه الگوهای معنایی (توهین جنسیتی، ظاهر) و نحوی (علائم تأکیدی، ساختار بدون عامل)، می‌توانند به عنوان نشانگرهای کلیدی برای توسعه یا بهبود

الگوریتم‌های تشخیص زورگیری سایبری بومی‌شده برای زبان فارسی به کار گرفته شوند. این الگوریتم‌ها باید قادر باشند فراتر از تطبیق کلیدواژه‌های ساده عمل کرده و الگوهای پیچیده‌تر را شناسایی کنند. شناخت الگوهای زبانی رایج در زورگیری سایبری می‌تواند مبنای طراحی برنامه‌های آموزشی مؤثرتری برای نوجوانان، جوانان، والدین و مربیان در زمینه سواد دیجیتال و مقابله با خشونت برخط قرار گیرد. آگاهی از این الگوها می‌تواند به کاربران کمک کند تا محتوای زورگیرانه را بهتر شناسایی کرده و نسبت به آن واکنش مناسب نشان دهند. شرکت‌های مالک شبکه‌های اجتماعی می‌توانند از این یافته‌ها برای بهبود سیستم‌های پایش محتوا و اعمال سیاست‌های خود در قبال کاربران فارسی‌زبان استفاده کنند تا با دقت و حساسیت فرهنگی بیشتری با محتوای آسیب‌زا مقابله نمایند.

پژوهش حاضر علی‌رغم تلاش برای رعایت دقت روش‌شناختی، با محدودیت‌هایی نیز روبرو بود. استفاده از نمونه‌گیری هدفمند و حجم نمونه ۵۰۰ موردی، اگرچه برای یک مطالعه کیفی مناسب است، اما قابلیت تعمیم یافته‌ها به کل جامعه کاربران ایرانی اینستاگرام و ایکس را محدود می‌کند. این پژوهش صرفاً بر تحلیل محتوای متنی تمرکز داشت و ابعاد بصری (تصاویر، ویدئوها، پانتومیم‌ها) که نقش مهمی در زورگیری سایبری (به‌ویژه در اینستاگرام) دارند، مورد بررسی قرار نگرفتند. زبان در فضای مجازی بسیار پویا و در حال تغییر است. الگوهای شناسایی‌شده ممکن است در طول زمان دستخوش تغییر شوند. تشخیص نیت زورگیرانه صرفاً بر اساس متن گاهی دشوار است و نیاز به درک بافت وسیع‌تری دارد که همیشه در دسترس نیست. تمرکز تنها بر روی اینستاگرام و ایکس، سایر سکوها محبوب مانند تلگرام یا سکوها داخلی را پوشش نداده است.

پیشنهادهای پژوهش

با توجه به یافته‌ها و محدودیت‌های این پژوهش، پیشنهادهای زیر برای تحقیقات آینده مطرح می‌شود:

گسترش حجم و تنوع نمونه: انجام پژوهش‌های مشابه با حجم نمونه بزرگتر و با استفاده از روش‌های نمونه‌گیری تصادفی‌تر (در صورت امکان) برای افزایش قابلیت تعمیم‌پذیری. تحلیل چندوجهی^۱: بررسی نقش عناصر بصری (تصاویر، ویدئوها، ایموجی‌ها) در کنار متن در زورگیری سایبری. مطالعات طولی: بررسی تغییرات الگوهای زبانی زورگیری سایبری در طول زمان.

1 multimodal analysis

مقایسه بین سکوی و درون‌زبانی: مقایسه دقیق‌تر الگوهای زبانی در سکوهای مختلف (شامل تلگرام، واتس‌اپ، سکوهای داخلی) و همچنین مقایسه بین گونه‌های مختلف زبان فارسی (مثلاً فارسی ایران و افغانستان).

توسعه و ارزیابی ابزارهای تشخیص: طراحی و پیاده‌سازی الگوریتم‌های تشخیص زورگیری سایبری بر اساس نشانگرهای زبانی شناسایی شده در این پژوهش و ارزیابی کارایی آن‌ها بر روی داده‌های واقعی.

بررسی جنبه‌های کاربردشناختی و گفتمانی: تحلیل عمیق‌تر نقش بافت، روابط بین کاربران، و راهبردهای‌های گفتمانی پیچیده‌تر در زورگیری سایبری

مطالعه قربانیان و شاهدان: بررسی تجربیات و برداشت‌های قربانیان و شاهدان زورگیری سایبری از زبان مورد استفاده و تأثیرات آن

زورگیری سایبری یک معضل جدی در فضای مجازی ایران است که نیازمند توجه و اقدام چندجانبه است. این پژوهش با تمرکز بر بُعد زبانی این پدیده، تلاش کرد تا درک عمیق‌تری از چگونگی نمود آن در زبان فارسی و در سکوهای پرکاربرد فراهم آورد. یافته‌ها بر اهمیت شناخت ویژگی‌های منحصربه‌فرد زبانی و فرهنگی در مقابله با این پدیده تأکید دارند و امید است که مبنایی برای تحقیقات آتی و اقدامات عملی مؤثرتر در راستای ایجاد فضای مجازی امن‌تر برای کاربران ایرانی باشد.

سیاسگزاری

از تمامی حامیان مادی و معنوی پژوهش که با حمایت‌هایشان زمینه‌های تحقق اهداف علمی را فراهم کردند، صمیمانه تقدیر و تشکر می‌نمایم.

منابع

زارع، حسین، علیمردادی مهتاب و رحمانیان، مهدیه. (۱۳۹۹). اثربخشی آموزش مهارت‌های زندگی بر وابستگی به فضای مجازی و قلدری سایبری در نوجوانان منطقه سه شهر تهران. *فصلنامه علمی پژوهش در یادگیری آموزشگاهی و مجازی*، ۸(۲)، ۹-۲۰. doi: 10.30473/etl.2020.50601.3147

فلاحی، وحید، نریمانی، محمد و عطادخت، اکبر. (۱۴۰۱). مدل معادلات ساختاری نقش باورهای غیرمنطقی در پیش‌بینی تجربه قلدری سایبری دانش‌آموزان با میانجیگری: دشواری در تنظیم هیجان. *مطالعات روان‌شناختی*، ۱۸(۱)، ۳۹-۵۳. doi: 10.22051/psy.2021.30477.2166

محمدی، فریبا، حسینیان، سیمین، یزدی، سیده منور، و مردانی راد، مژگان. (۱۴۰۱). نقش میانجی‌گری خشم در رابطه بین حمایت اجتماعی ادراک شده با قلدری-قربانی سایبری در دختران نوجوان. *فصلنامه فرهنگی-تربیتی زنان و خانواده*، ۱۷(۶۰)، ۱۳۱-۱۵۳.

Austin, J. L. (1962). *How to do things with words*. Oxford University Press.

- Badri, M., Nuaimi, A. A., Guangul, F. M., & Rashedi, A. (2021). Social media usage and challenges: A study in the United Arab Emirates. *International Journal of Information Management Data Insights*, 1(2), 100038.
- Baker, P. (2006). Using corpora in discourse analysis. Continuum. Braun, V., & Clarke, V. (2006). Using thematic analysis in psychology. *Qualitative Research in Psychology*, 3(2), 77–101.
- Bauman, S. (2013). Cyberbullying: What Does Research Tell Us? *Theory Into Practice*, 52(4), 249–256. <https://doi.org/10.1080/00405841.2013.829727>
- Braun, V., & Clarke, V. (2006). Using thematic analysis in psychology. *Qualitative Research in Psychology*, 3(2), 77–101.
- Brewer, G., & Kerslake, J. (2015). Cyberbullying, self-esteem, empathy and loneliness. *Computers in Human Behavior*, 48, 255–260.
- Lonigro, A., Schneider, B. H., Laghi, F., Baiocco, R., Pallini, S., & Brunner, T. (2015). Is cyberbullying related to trait or state anger?. *Child Psychiatry & Human Development*, 46, 445–454.
- Dadvar, M., Trieschnigg, D., Ordelman, R., & de Jong, F. (2013). Improving cyberbullying detection with user context. In *Advances in Information Retrieval: 35th European Conference on IR Research, ECIR 2013* (pp. 693–696). Springer.
- Dinakar, K., Jones, B., Havasi, C., Lieberman, H., & Picard, R. (2012). Common sense reasoning for detection, prevention, and mitigation of cyberbullying. *ACM Transactions on Interactive Intelligent Systems (TiiS)*, 2(3), 1–30.
- Festl, R., Scharrow, M., & Quandt, T. (2013). Peer influence, internet use and cyberbullying: A comparison of different context effects among German adolescents. *Journal of Children and Media*, 7(4), 486–501.
- Hinduja, S., & Patchin, J. W. (2010). Bullying, cyberbullying, and suicide. *Archives of Suicide Research*, 14(3), 206–221.
- Hinduja, S., & Patchin, J. W. (2015). *Bullying beyond the schoolyard: Preventing and responding to cyberbullying* (2nd ed.). Corwin Press.
- Kaplan, A. M., & Haenlein, M. (2010). Users of the world, unite! The challenges and opportunities of social media. *Business Horizons*, 53(1), 59–68. <https://doi.org/10.1016/j.bushor.2009.09.003>
- Kowalski, R. M., Giumetti, G. W., Schroeder, A. N., & Lattanner, M. R. (2014). Bullying in the digital age: A critical review and meta-analysis of cyberbullying research among youth. *Psychological Bulletin*, 140(4), 1073–1137. [doi: 10.1037/a0035618](https://doi.org/10.1037/a0035618).
- McEnery, T., & Wilson, A. (2001). *Corpus linguistics: An introduction* (2nd ed.). Edinburgh University Press.
- Nahar, V., Li, X., & Pang, C. (2012). An effective approach for cyberbullying detection. *Communications in Information Science and Management Engineering*, 2(9), 238.
- Olweus, D. (2012). Cyberbullying: An overrated phenomenon? *European Journal of Developmental Psychology*, 9(5), 520–538.

- Patchin, J. W., & Hinduja, S. (2012). *Cyberbullying: An update and synthesis of the research*. In *Cyberbullying prevention and response: Expert perspectives* (pp. 13-35). Routledge.
- Patton, M. Q. (2015). *Qualitative research & evaluation methods: Integrating theory and practice (4th ed.)*. Sage publications.
- Reynolds, K., Kontostathis, A., & Edwards, L. (2011). Using machine learning to detect cyberbullying. In *Proceedings of the 10th International Conference on Machine Learning and Applications Workshops* (Vol. 2, pp. 241-244). IEEE.
- Runions, K., Shapka, J. D., Dooley, J., & Modecki, K. (2013). Cyber-aggression and victimization and social information processing: Integrating the medium and the message. *Psychology of Violence*, 3(1), 9-26. <https://doi.org/10.1037/a0030511>
- Salawu, S., He, Y., Lumsden, J., & Olatunji, S. O. (2020). Approaches towards the detection of cyberbullying: A survey. *IEEE Transactions on Affective Computing*, 13(2), 861-880.
- Searle, J. R. (1969). *Speech acts: An essay in the philosophy of language*. Cambridge University Press.
- Searle, J. R. (1975). *A taxonomy of illocutionary acts*. In *Language, mind, and knowledge* (pp. 344-369). University of Minnesota Press.
- Slonje, R., & Smith, P. K. (2008). Cyberbullying: Another main type of bullying? *Scandinavian Journal of Psychology*, 49(2), 147-154.
- Smith, P. K., López-Castro, L., Robinson, S., & Görzig, A. (2020). Consistency of findings on cyberbullying: A comparison of two national studies in England. *Nordic Psychology*, 72(1), 6-19.
- Smith, P. K., Mahdavi, J., Carvalho, M., Fisher, S., Russell, S., & Tippett, N. (2008). Cyberbullying: Its nature and impact in secondary school pupils. *Journal of Child Psychology and Psychiatry*, 49(4), 376-385.
- Tokunaga, R. S. (2010). Following you home from school: A critical review and synthesis of research on cyberbullying victimization. *Computers in Human Behavior*, 26(3), 277-287.
- Van Hee, C., Lefever, E., Hoste, V., Jacobs, G., & Montero Perez, M. (2015). Detection of cyberbullying in social media text. In *Proceedings of the 6th Workshop on Computational Approaches to Subjectivity, Sentiment and Social Media Analysis* (pp. 31-35).
- Waseem, Z., & Hovy, D. (2016). Hateful symbols or hateful people? predictive features for hate speech detection on twitter. In *Proceedings of the NAACL student research workshop* (pp. 88-93).