

Exposure To AI Content and The Rationalisation of Digital Privacy Abuse: A Study Among Students at University X in Pekanbaru

Sa'diyah Aristawati Septini¹⁾, Fakhri Usmita²⁾

^{1,2} Islamic University of Riau, Indonesia

Email:sadiyaharistawatiseptini@student.uir.ac.id¹, fakhri@soc.uir.ac.id²

Abstract: *The rapid development of generative Artificial Intelligence (AI) has triggered the "photos with idols" Prompt-AI content trend, which manipulates visual identities on social media,. This study aims to measurably analyze the relationship between exposure to such Prompt-AI content and the tendency to neutralize (normalize) digital privacy abuse among students at University X in Pekanbaru. The research employed a quantitative approach with a descriptive design. Data were collected via online questionnaires using a four-point Likert scale from 103 student respondents selected through purposive sampling,. Data analysis was conducted using the Pearson Product Moment correlation test via SPSS software, following validity and reliability confirmation (Cronbach's Alpha > 0.80) and ensuring normal data distribution through the One-Sample Kolmogorov-Smirnov test,. The results demonstrate a significant positive relationship of moderate strength ($r=0.535$; $p=0.000$) between content exposure and the neutralization of digital privacy abuse.*

Abstrak: *Pesatnya perkembangan AI generatif memicu tren konten Prompt-AI "foto bersama idola" yang memanipulasi identitas visual di media sosial,. Penelitian ini bertujuan untuk menganalisis secara terukur hubungan antara tingkat terpapar konten Prompt-AI tersebut dengan kecenderungan netralisasi penyalahgunaan privasi digital pada mahasiswa Universitas X di Pekanbaru. Metode penelitian menggunakan pendekatan kuantitatif dengan desain deskriptif. Data dikumpulkan melalui kuesioner daring menggunakan skala Likert empat poin terhadap 103 responden mahasiswa yang dipilih melalui teknik purposive sampling. Analisis data dilakukan menggunakan uji korelasi Pearson Product Moment melalui perangkat lunak SPSS, setelah instrumen dinyatakan valid dan reliabel (Cronbach's Alpha > 0,80) serta data terdistribusi normal melalui uji One-Sample Kolmogorov-Smirnov. Hasil penelitian menunjukkan hubungan positif yang signifikan dengan kekuatan sedang ($r=0,535$; $p=0,000$) antara terpapar konten dengan netralisasi penyalahgunaan privasi digital.*

Keywords: Artificial Intelligence, Neutralisation, Digital Privacy

INTRODUCTION

The development of artificial intelligence (AI) technology has brought about significant changes in various digital activities, particularly through generative AI innovations capable of producing visual, textual and audio content via text instructions known as prompts (Dang et al., 2022). The ease of use of this technology has made it extremely popular amongst students—as part of Generation Z—who regard it as practical and efficient (Khemiko et al., 2024). However, the intensity of its use raises serious digital ethical issues, particularly when AI technology is used to manipulate a person's visual identity without

the owner's consent. This phenomenon is evident in the trend of 'photos with idols' generated by Prompt-AI, which is widely circulating on social media platforms such as TikTok and X, where users create realistic depictions of personal closeness with public figures (Koreaboo, 2024). Participation in this trend is often driven by the 'Fear of Missing Out' (FOMO) phenomenon, namely the desire to keep up with developments and trends that are currently popular on social media (Fahsyah & Junaidi, 2025).

From a legal perspective in Indonesia, the use of personal data including facial images without consent is regulated under Law No. 27 of 2022 on the Protection of Personal Data. The main issue arises when students, as an educated group, tend to view this manipulation as a form of acceptable entertainment, thereby blurring ethical boundaries (Kompas, 2025). From a criminological perspective, this tendency to justify such behaviour can be explained through Sykes and Matza's theory of neutralisation, whereby individuals rationalise deviant behaviour in order to maintain a moral self-identity (Chaerunnisa & Hakim, 2025). Previous research has extensively examined the misuse of AI in the context of deepfake pornography or legal aspects (Harun & Nurhadiyanto, 2024; Kurniarullah et al., 2024; Sembiring & Firdaus, 2025). However, studies that empirically analyse the relationship between exposure to digital trends and neutralisation mechanisms amongst university students remain limited.

This study aims to fill this gap by being conducted at University X in Pekanbaru, an institution grounded in Islamic values and Malay culture that upholds ethics and individual dignity. The significance of this study lies in its attempt to measure the extent to which the intensity of exposure to digital trends can erode students' ethical sensitivity towards others' privacy. Based on this issue, this study aims to analyse the relationship between exposure to Prompt-AI content featuring 'photos with idols' and the tendency towards the normalisation of digital privacy violations among students at University X in Pekanbaru.

Content exposure is understood as the process by which an individual receives message material via digital platforms, triggering the emergence of attention and the formation of attitudes (Dini, 2025). In the context of generative technology, the level of students' exposure to "photos with idols" content is measured through three main indicators: frequency, duration and attention (Ikhsan & Daherman, 2022). Frequency refers to the regularity of exposure, whilst duration encompasses the length of time spent consuming the content. Attention is a crucial element indicating students' cognitive engagement, where their focus is often divided between interest in the visual output and the technical instructions or prompts used within the AI system.

The neutralisation mechanism is a socio-psychological process whereby individuals rationalise or justify behaviour that violates social norms (Sykes & Matza, 2017). In the practice of visual identity manipulation in cyberspace, there are two main forms of justification used to suspend moral conflict:

- 1) Denial of Injury: Individuals tend to assess that the manipulation of an idol's photograph does not pose any real danger or cause serious harm to the public figure in question. Such content is positioned solely as 'safe' static digital entertainment.
- 2) Denial of Victim: The perception arises that public figures are not genuine victims because their status as celebrities is deemed to inherently erode the boundaries of absolute privacy (Sahan et al., 2025). This loss of perception regarding the victim serves as a strong basis for students to justify tolerating the unauthorised use of others' images.

METHOD

This study employs a quantitative design with a descriptive approach to examine the phenomenon objectively (Muin, 2023). The population is categorised as an infinite population (Saptohadhi, 2022), namely all students who have ever been exposed to AI-based visual manipulation content on social media. The sample was selected explicitly using purposive sampling with the following main criteria: (1) being an active student, and (2) having experience of exposure to Prompt-AI content featuring 'photos with idols' on social media. The sample size was set at 103 respondents, calculated using the Lemeshow formula (Agustin & Mindiharto, 2025) with a margin of error of 10 per cent, to which was increased by 7% as a buffer to ensure data stability and to anticipate the presence of invalid instruments during the data processing stage.

Primary data were collected via an online questionnaire using Google Forms, employing a four-point Likert scale to encourage respondents to express their positions clearly (Mulyani, 2022). In addition to the questionnaire, limited observation of interactions on social media was also conducted to strengthen the understanding of patterns of justification for digital behaviour (Mulyani, 2022). All data were systematically analysed using SPSS software (Nurcahyati, 2023). The testing phase began with a Pearson Product Moment validity test to ensure the validity of the instrument's items (Khairunnisa, 2023) and a Cronbach's Alpha test to measure the reliability or consistency of the instrument's results. Subsequently, the One-Sample Kolmogorov-Smirnov test was used to assess the normality of the data distribution as a basis for selecting the appropriate statistical technique (Pujianto & Mulyati, 2022). As a final step to analyse the existence, direction and strength of the relationship between the intensity of exposure to digital trends and the mechanisms for neutralising the misuse of digital privacy, a Pearson Product Moment correlation test was conducted (Vikaliana, 2022).

RESULTS AND DISCUSSION

This study involved 103 students selected using purposive sampling. Based on the respondents' characteristics, the majority were female (65 per cent) and aged between 21 and 22 years (63.1 per cent). These characteristics indicate that the research sample was dominated by Generation Z students, who are active users of social media and digital technology.

Prior to hypothesis testing, the research instrument underwent validity and reliability testing. The results showed that all statement items were deemed valid as they had a calculated r value greater than the table r value (0.196). The reliability test also showed a good level of consistency, with a Cronbach's Alpha value of 0.814 for the Prompt-AI Content Exposure variable (X) and 0.908 for the Neutralisation of Digital Privacy Misuse variable (Y). Furthermore, the normality test using the One-Sample Kolmogorov-Smirnov method yielded a significance value of 0.200 ($p > 0.05$), meaning the data were deemed to be normally distributed and met the criteria for analysis using parametric statistics.

Table 1 Descriptive Summary of Prompt-AI Content Exposure and Neutralisation Mechanisms

Variable	Indicator	Main Trends	N	%
Exposure to Prompt-AI Content	Frequency (30 days)	Exposed 3–5 times in the last month	42	40,8
	Duration	Accessed content for 5–10 seconds	55	53,4
	Attention	Interested in reading the prompt or AI instructions	66	64,1
Neutralising the Misuse of Digital Privacy	Denial of Injury	Considering the impact to be minor provided it is not taken to excess	65	63,1
	Denial of Victim	Believing that public figures are better equipped to deal with the use of their image in the digital sphere	68	66,0

Source: Researcher's Data Analysis (2026)

Based on Table 1, the majority of respondents were exposed to 'photo with an idol' Prompt-AI content 3–5 times in the past month. Although the duration of exposure tended to be brief, respondents demonstrated a fairly high level of attention to the technical aspects of the content, particularly the content of the prompt used. Regarding the neutralisation variable, the most dominant forms were 'denial of victim' and 'denial of injury', indicating a tendency amongst respondents to view the use of public figures' images as an action that does not cause serious consequences.

Table 2 Results of the Pearson Product-Moment Correlation Test

Variable	r	Sig. (2-tailed)	N
Exposure to Prompt-AI Content (X) and Mitigation of Digital Privacy Abuse (Y)	0,535	0,000	103

Source: Researcher's Data Analysis (2026)

Based on Table 2, a correlation coefficient of 0.535 was obtained with a significance value of 0.000 ($p < 0.05$). These results indicate a positive and significant relationship between exposure to Prompt-AI content titled 'photos with idols' and the normalisation of digital privacy abuse amongst students. A correlation value of 0.535 falls within the category of a moderate relationship, meaning that the higher the level of exposure to such content, the greater the tendency for students to rationalise practices involving the abuse of digital privacy.

The intensity of exposure to Prompt-AI content indicates that visual manipulation has become part of students' daily media consumption routines. The characteristics of Generation Z, who have a high affinity for artificial intelligence systems, support the acceleration of this trend as it is perceived as more practical and efficient (Khemiko et al., 2024). However, students' high focus on the technical aspects of image generation, such as prompt text, often diverts their attention from the ethical implications regarding the manipulated subjects. From a criminological perspective, this situation triggers a situational moral shift or 'drift', whereby individuals tolerate deviant behaviour driven by the urge to remain socially relevant through the phenomenon of Fear of Missing Out (FOMO) (Fahsyah & Junaidi, 2025).

The process of moral justification for the unauthorised use of visual identities can be explained through the Neutralisation Techniques Theory (Sykes & Matza, 2017). The use of the 'denial of injury' technique indicates that students engage in moral reduction by interpreting photo manipulation as a 'safe' visual activity or merely personal entertainment that has no real impact on the private lives of public figures (Yasmine & Gusnita, 2024).

Furthermore, the prevalence of the 'denial of victim' technique reveals a pattern of objectification in which an idol's popularity is used as a justification for disregarding privacy boundaries. Students tend to view public figures as being professionally better equipped to handle the use of their images online; consequently, the loss of control over personal data is seen as a natural consequence of fame (Herdian & Sumarwan, 2025). This demonstrates that high technical digital literacy does not always correlate directly with ethical awareness regarding the protection of individual dignity in cyberspace (Lestari et al., 2026).

From an Islamic perspective, this phenomenon of neutralisation represents a form of unawareness of the harm caused, even though the perpetrator perceives their actions as perfectly normal, as highlighted in Surah Al-Baqarah, verses 11–12. An academic environment grounded in Islamic values should emphasise the principle of *hifz al-'irdh* (preserving honour) as the primary boundary in interacting with technology (Muin, 2023). The principle of moral responsibility in Surah Al-Isra' verse 36 affirms that every digital activity including what is viewed and observed on social media—will be

held to account before Allah SWT. Therefore, strengthening moral literacy is crucial so that students possess the integrity to safeguard privacy and human dignity as subjects in the digital space.

CONCLUSION

This study reveals a significant positive correlation of moderate strength ($r = 0.535$) between the intensity of exposure to generative AI-based visual manipulation content and the tendency to rationalise digital privacy violations amongst the student group who were the subjects of the study. These findings indicate that the frequency, duration and attention paid to digital trends facilitate the use of neutralisation techniques as a socio-psychological mechanism to minimise moral conflict when interacting with manipulative content. Students tend to apply the ‘denial of injury’ technique by perceiving visual manipulation as entertainment material that does not cause any tangible material or reputational harm to the subject. Furthermore the ‘denial of victim’ mechanism emerges through the objectification of public figures, where popularity is often seen as justifying a reduction in an individual’s privacy in the digital sphere. This phenomenon reflects a shift in situational moral drift, indicating a disconnect between digital usage and an awareness of ethical responsibility in upholding human dignity. As a practical implication for educational institutions, there is a need to strengthen moral literacy that focuses not only on technological proficiency but also on an understanding of legal boundaries in accordance with Law No. 27 of 2022 on the Protection of Personal Data. Social media platform developers are advised to optimise content moderation through the automatic labelling of AI-generated synthetic content to prevent the normalisation of the misuse of visual identities in the public sphere. This study is limited by the homogeneity of the population within a single academic region; therefore, generalising the findings to a broader social context must be done with caution. Consequently, future researchers are recommended to explore deeper socio-psychological motivations through a mixed-methods approach to gain a more comprehensive understanding of the dynamics of generative digital culture.

REFERENCES

- Agustin, M. A. S., & Mindiharto, S. (2025). Faktor-Faktor Yang Berhubungan Dengan Tingkat Kelelahan Kerja Pada Pengemudi Ojek Online Di Gresik Tahun 2025. *Jurnal Mutiara Kesehatan Masyarakat*, 10(1), 43-54.
- Chaerunnisa, A. H., & Hakim, F. N. (2025). Analisis Teori Teknik Netralisasi Pada Tindakan Penyalahgunaan Obat Lambung Sebagai Alat Aborsi Yang Diperjualbelikan Di Facebook. *Ranah Research: Journal Of Multidisciplinary Research And Development*, 7(5), 3862-3874.
- Dang, H., Mecke, L., Lehmann, F., Goller, S., & Buschek, D. (2022). How To Prompt? Opportunities And Challenges Of Zero-And Few-Shot Learning For Human-Ai Interaction In Creative Applications Of Generative Models. *Arxiv Preprint Arxiv:2209.01390*

- Dini, M. R. (2025). Pengaruh Terpaan Konten Instagram Penggunaan Aplikasi Shopee Melalui Instagram Shopee_Idl. *Journal Of Scientific Communication (Jsc)*, 7(1).
- Fahsyah, N. K. P. A., & Junaidi, A. (2025). Fenomena Fear Of Missing Out (Fomo) Pada Generasi Z Dalam Mengikuti Trend Tiktok. *Kiwari*, 4(1), 61- 70.
- Harun, F. A., & Nurhadiyanto, L. (2024). Rekayasa Konten Pornografi Berbasis AI Image Generator Dalam Perspektif Space Transition Theory. *Ranah Research: Journal Of Multidisciplinary Research And Development*, 6(3), 408-418
- Herdian, A., & Sumarwan, U. (2025). Analisis Kriminologi Deepfake Melalui Media Sosial Berdasarkan Teori Rational Choice. *IKRA-ITH HUMANIORA: Jurnal Sosial Dan Humaniora*, 9(1), 323-331
- Ikhsan, M., & Daherman, Y. (2022). Pengaruh Terpaan Tayangan Review Gadget Di Youtube Terhadap Minat Beli Anggota Komunitas Game@ Freefireriaul. *Medium: Jurnal Ilmiah Fakultas Ilmu Komunikasi Islam Riau*, 10(02), 1-11.
- Khairunnisa. (2023). Uji Validitas Dan Reliabilitas. Dalam Buku Ajar Metode Penelitian (. 53 - 68). Pangkal Pinang : CV Science Techno Direct
- Khemiko, K., Kornarius, Y. P., Caroline, A., Gusti, T. E. P., & Gunawan, A. (2024). IPengaruh Sikap Generasi Milenial Dan Generasi Z Terhadap Kecenderungan Untuk Terus Menggunakan Teknologi Kecerdasan Buatan. *Ekonomika45: Jurnal Ilmiah Manajemen, Ekonomi Bisnis, Kewirausahaan*, 12(1), 338-353.
- Koreaboo. (2025, September 11). Dangerous “sexualized” AI-created trend with male idols sparks major disgust.
- Kurniarullah, M. R., Nabila, T., Khalidy, A., Tan, V. J., & Widiyani, H. (2024). Tinjauan Kriminologi Terhadap Penyalahgunaan Artificial Intelligence: Deepfake Pornografi Dan Pencurian Data Pribadi. *Jurnal Ilmiah Wahana Pendidikan*, 10(10), 534-547
- Lestari, P. A., Saputri, Z. R., Jamilah, R., Febiyanti, F., Fadilah, M. F., Firnanda, M. F., & Koeswara, T. S. N. (2026). Konten Ilegal Kecerdasan Buatan Berbasis Deepfake dan Upaya Penanggulangannya di Indonesia. *Jurnal Ilmiah Sistem Informasi*, 5(1), 26-35
- Muin, A. (2023). Metode Penelitian Kuantitatif. Malang : Cv. Literasi Nusantara Abadi.
- Mulyani, Wiwiek. (2022). Bab X Teknik Pengumpulan Data. Dalam Hidayanti Dan Aulia (Edisi I), *Metodologi Penelitian Kuantitatif* (. 116 - 131). Padang
- Novia Respati. (2025, September 22). Tren foto bareng idola pakai AI, apa yang dicari? *Katanetizen Kompas*.
- Nurchayati, Sri. (2023). Analisis Data Epidemiologis. Dalam Buku Ajar Metode Penelitian (. 129 - 141). Pangkal Pinang : Cv Science Techno Direct
- Pujianto, Aging Dan Mulyati, Awin. (2022) Uji Asumsi Klasik. Dalam Surur, Miftahus (Edisi I) *Ragam Penelitian Dengan SPSS* (15 – 31). Sukoharjo ; CV Tahta Media Group

- Sabtohadhi, Joko. (2022). Bab Viii Populasi, Sampel Dan Variabel Penelitian. Dalam Hidayanti Dan Aulia (Edisi I), Metodologi Penelitian Kuantitatif (. 79 - 91). Padang : PT. Global Eksekutif Teknologi
- Sahan, M. Y., Gual, Y. A., & Bello, M. F. Y. (2025). Eksplorasi Manajemen Privasi Komunikasi Mahasiswa Melalui Buku Harian Dan Instagraml. Wacana: Jurnal Ilmiah Ilmu Komunikasi, 24(1), 178-189.
- Sembiring, J., & Firdaus, S. U. (2025). Criminal Liability For Misuse Of Artificial Intelligence (Ai) In Deepfake Crimes”. International Journal Of Multi Science, 5(01), 1-9.
- Sykes, G. M., & Matza, D. (2017). Techniques Of Neutralization: A Theory Of Delinquency. In Delinquency And Drift Revisited, Volume 21 (Pp. 33-41). Routledge.
- Vikaliana, Resista. (2022) Analisa Regresi. Dalam Surur, Miftahus (Edisi I) Ragam Penelitian Dengan SPSS (1 – 14). Sukoharjo ; CV Tahta Media Group
- Yasmine, A. F., & Gusnita, C. (2024). Fenomena Jokes Seksis Mahasiswa Sebagai Bentuk Normalisasi Pelecehan Seksual Secara Verbal. Ranah Research: Journal of Multidisciplinary Research and Development, 6(4), 528-536.