

Convergent Parallel Mixed-Method Study of Frameworks and Consideration of Ethical AI-Assisted Research Paper

Juvenal S. Dalimeg¹, Antonio Ponciano L. Rodil²,
Kimberly Ann S. Cantilero³, Michael U. Jimenez⁴,
Jan Ashley D. Casco⁵, Kezia Rapha de Guzman⁶

¹ National University, Philippines; National University Clark, Mabalacat City Philippines (Psychology/Faculty/School of Education, Arts and Sciences/ School, University), Philippines, juvenaldalimeg@gmail.com, 0009-0001-0791-4337

² National University, Philippines; National University Clark, Mabalacat City Philippines (Psychology/Faculty/School of Education, Arts and Sciences/School, University), Philippines, alrodil@nu-clark.edu.ph, 0000-0002-4363-1246

³ National University, Philippines; National University Clark, Mabalacat City Philippines (Psychology/Faculty/School of Education, Arts, and Sciences/School, University), Philippines, kscantilero@nu-clark.edu.ph, 0000-0002-3642-859X

⁴ National University, Philippines; National University Clark, Mabalacat City Philippines (General Education/Faculty/School of Education, Arts, and Sciences/School, University), Philippines, jdcasco@nu-clark.edu.ph, 0009-0007-8355-5034

⁵ National University, Philippines; National University Clark, Mabalacat City Philippines (Psychology/Faculty/School of Education, Arts, and Sciences/School, University), Philippines, 0009-0004-0941-8232

Artificial intelligence (AI) and Natural Language Processing (NLP) technologies are being adopted both in academic research work and the emerging development of academic writing. The benefits these tools provide, however, come with ethical issues that challenge authorship, transparency, and research integrity. To understand the behavioral AI-interactions of academic stakeholders, this study employs Cognitive Appraisal Theory and the framework of moral distress to explore the

*Author info: Correspondence should be sent to: Juvenal Dalimeg,
Email: jsdalimeg@nu-clark.edu.ph*

North American Journal of Psychology, 2026, Vol. 28, No. 3, 1932-1955
© NAJP <https://doi.org/10.65696/001c.168759>

psychological affordances and constraints of AI-assisted writing tools. Utilizing a convergent parallel mixed-methods design, data were gathered from 200 participants, including academic researchers, graduate students, university administrators, AI developers, and ethicists. The results indicate that cognitive appraisals of threat and benefit reliably predict ethical concerns, moderated significantly by professional role and demographic group. This study outlines how policymakers, researchers, and institutions can mitigate moral distress and ethically integrate AI into academic writing.

Keywords: Artificial Intelligence (AI), Natural Language Processing (NLP), Cognitive Appraisal Theory, Moral Distress, Academic Dishonesty, Ethical Framework

INTRODUCTION

Artificial intelligence and natural language processing (AI/NLP) technologies have rapidly transformed academic research writing worldwide. These advancements have introduced tools such as ChatGPT and GPT-4 (OpenAI, 2023), which are capable of generating coherent, grammatically proficient text that assists scholars in drafting research manuscripts. However, the integration of AI into academic writing presents substantial concerns, including questions of authorship, process transparency, algorithmic bias, and broader research ethics (Bender et al., 2021). In the Philippines, for example, hybrid analog-digital systems of AI in education are still nascent, and there are no frameworks for the ethical use of AI yet.

With the sudden expansion of their educational sector and an increasing use of technology in both westernized economies as well as underdeveloped countries, such as the Philippines, we have a particularly layered argument for examining the ethical concerns surrounding AI-led research paper production. Despite the progress made to digitally empower education in the Philippines — most notably during the national COVID-19 response where there was a rapid move online and digital adoption (UNESCO, 2021), there are still challenges (Alampay, 2020) including the absence of sufficient technology access and usage, the need for more digital literacy as well as policy definition on the use of AI in educational institutions.

In recent years, artificial intelligence has been recognized by the Philippine government as a critical lever for driving economic growth and innovation. The National AI Roadmap 2021, created by the Department of Science and Technology (DOST), integrates AI across a wide range of sectors, including education (DOST, 2021). However, AI adoption among universities in the Philippines remains patchy (Magno & Serrano, 2022), with top-tier institutions acting as trailblazers while under-resourced schools gradually work to catch up.

AI tools are mainly used in Philippine higher education for administrative purposes like enrolling students and doing performance data analyses (Tongson et al., 2023). But AI's potential to help academic writing and research has not yet been realized by many. According to a study done at De La Salle University, AI tools are increasingly being used by faculty and students for drafting and editing, but awareness of the ethical considerations involved in using them is low (Santos et al., 2022).

The ethical quandary raised in the context of AI-assisted research paper generation is acute in the Philippines (Bernardo et al., 2021), as academic integrity and plagiarism are already an issue here. Such problems are further compounded by the absence of robust guidelines on the use of AI tools for academic writing, which, in turn, leaves institutions and researchers to work out these issues on their own. Moreover, the potential to compound the disparate and stratified education system through AI tools presents further concerns of its own (Magno & Serrano, 2022).

To understand the behavioral AI-interactions of academic stakeholders, this study employs Cognitive Appraisal Theory to examine how individuals evaluate the integration of AI tools in research. The perceived benefits and perceived limitations of AI-generated content function as primary cognitive appraisals. When stakeholders perceive high threat, it triggers moral distress—a psychological phenomenon occurring when one recognizes the ethically appropriate action but feels constrained by institutional gaps or technological opacity. By investigating these cognitive biases and threat perceptions, this study elucidates the psychological underpinnings of ethical decision-making in human-AI collaboration.

METHODS

Research Design

This study employed a Convergent Parallel Mixed-Methods Design to fully examine the Ethical Considerations surrounding the use of Artificial Intelligence-Assisted Research within the Academic Environment in the Philippines. A Convergent Parallel Design was used to collect and analyze both quantitative and qualitative data at the same time. By using this type of design, both methods can be combined to create a better understanding of the Phenomenon than each method could have created separately (Creswell & Plano Clark, 2018).

Quantitative Strand

A Descriptive-Correlational Cross-Sectional Study Methodology was used as part of the quantitative strand. Using a structured survey instrument, a total of 200 participants (academic researchers, graduate students, university administrators, AI developers, and research ethics experts) provided responses for gathering the data. In addition to collecting information about perceived risk and benefit associated with the use of AI-

Assisted Writing, inferential statistics (one-way ANOVA and multiple regression) will be used to determine which of those variables are most predictive of an individual's level of concern regarding ethics.

Qualitative Strand

At the same time, a qualitative phenomenological strand was also implemented by conducting semi-structured interviews with thirty key informants. This strand seeks to gain insight into the unique perspectives of various stakeholder groups toward issues related to transparency, bias, and academic integrity. Because the qualitative strand provides a greater level of detail than standardized scales regarding the moral reasoning and institutional obstacles experienced by the stakeholders interviewed in this study, it serves to add depth to the statistically generated findings.

Integration and Data Triangulation

As specified under the guidelines of the Convergent Parallel Design, the two strands of data collected throughout this study will be analyzed individually and then during the interpretation process they will be merged together. While the quantitative results will provide generalizable trends concerning demographics, the qualitative themes will help support, explicate, or contradict those quantitative results. Through this triangulation process, a more comprehensive and empirically based framework for developing ethical standards for the use of AI-Assisted Writing will be developed.

Participants

The present study applied a stratified and purposeful sampling method to provide a broad, representative sample of stakeholder groups in the academic and technological environments of the Philippines. From universities, organizations that develop artificial intelligence, and institutional review boards, a total of 200 individuals ($N = 200$) participated in this study. In terms of their membership to one or more categories, the sample was composed of: Academic Researchers and Faculty ($n = 80$), i.e., those who are full-time employees at an institution and have produced articles for scholarly publication; Graduate Students ($n = 60$); students pursuing either master's or doctoral degrees at an institution of higher education and having some experience using digital research tools; University Administrators ($n = 30$), e.g., high-ranking administrative personnel whose responsibility includes making university policies and overseeing academic integrity; AI/NLP Developers ($n = 20$); professionals who specialize in creating, developing, implementing, and deploying natural language processing models; Research Ethics Specialists ($n = 10$); IRB members or other specialists who oversee the ethical implications of emerging technologies.

Instruments

The study utilized a dual-component instrumentation approach to

facilitate the convergent parallel design, consisting of a structured survey for quantitative data and a semi-structured interview guide for qualitative insights.

1. Quantitative Survey Instrument

The survey was divided into three distinct sections, designed to capture the variables outlined in the theoretical framework:

Section A: Demographic and Professional Profile. This section collected baseline data, including age, gender, institutional affiliation, academic rank, and the participants' self-reported level of technical proficiency with AI and Natural Language Processing (NLP) tools.

Section B: Perceptions of AI-Assisted Writing. This 10-item scale measured two primary constructs: Perceived Benefits (e.g., efficiency, linguistic enhancement) and Perceived Risks (e.g., erosion of critical thinking, dependency). Responses were recorded on a 5-point Likert scale, ranging from 1 (Strongly Disagree) to 5 (Strongly Agree).

Section C: Ethical Concerns Scale. This section evaluated participants' levels of concern regarding four pillars of research ethics: authorship, transparency, reproducibility, and research integrity.

Reliability Analysis: To ensure the internal consistency of the quantitative instrument, a pilot test was conducted. The Cronbach's Alpha coefficients were calculated as follows:

Perceptions of AI Scale (Section B): $\alpha = 0.88$.

Ethical Concerns Scale (Section C): $\alpha = 0.92$.

While internal consistency was established via Cronbach's alpha ($\alpha = 0.88$ and $\alpha = 0.92$), content validity was ensured through a rigorous review by a panel of three experts in research ethics and educational psychology. Furthermore, the construct validity of the Perceptions of AI Scale and Ethical Concerns Scale was evaluated utilizing an Exploratory Factor Analysis (EFA).

Prior to extraction, the suitability of the data for factor analysis was assessed. The Kaiser-Meyer-Olkin (KMO) measure verified the sampling adequacy for the analysis, $KMO = 0.86$, which is well above the recommended threshold of 0.60. Bartlett's Test of Sphericity indicated that the correlation matrix was not an identity matrix, $\chi^2(45) = 842.15$, $p < .001$, confirming that the item correlations were sufficiently robust for EFA.

Given the theoretical assumption that cognitive appraisals of AI are interconnected, a Principal Axis Factoring (PAF) extraction method with a Promax (oblique) rotation was applied. The analysis yielded a distinct two-factor solution based on the Kaiser criterion (eigenvalues > 1.0), which cumulatively explained 64.3% of the total variance.

Factor 1 (Perceived Threat/Risk): Accounted for 38.1% of the initial variance (eigenvalue = 4.12). Item loadings for this factor ranged from 0.65 to 0.88, capturing constructs such as algorithmic bias, lack of

transparency, and the erosion of academic integrity.

Factor 2 (Cognitive Benefit): Accounted for 26.2% of the variance (eigenvalue = 2.85). Item loadings ranged from 0.61 to 0.82, encompassing elements like research efficiency and linguistic enhancement.

All items demonstrated strong primary loadings (> 0.50) with no problematic cross-loadings (secondary loadings < 0.30), thereby confirming the distinct psychometric dimensionality and construct validity of the instruments used in this study.

2. A qualitative semi-structured interview guide was developed to elicit deep, narrative-driven data from a subset of participants ($n = 30$). The guide was validated by a panel of three experts in ethics and educational technology. Key inquiry areas included:

3. Personal experiences with AI-mediated authorship.
4. Institutional gaps in AI policy and disclosure requirements.
5. Moral dilemmas encountered when utilizing NLP tools for data interpretation.
6. Open-ended prompts allowed participants to elaborate on the "why" behind the quantitative trends observed in the survey phase.

Procedure

Phase 1 — Ethical Approval and Participant Selection

Before we began collecting data for this study, our research proposal was formally reviewed and approved through the Institutional Review Board (IRB) at [institution name]. Once IRB approval had been received, participants were selected through purposive sampling. To be eligible to participate in this study, participants needed to have demonstrated experience as either an academic publisher; professional AI developer; or a member of one of their institution's research ethics committees. Each participant also signed a formal informed consent form that detailed the goals of the study, described how all aspects of the study would remain confidential, explained how the participant could withdraw from participating in the study at any time without penalty, and stated that participation in the study was completely voluntary.

Phase 2 – Quantitative Data Collection (web survey)

To collect quantitative data, we created a web-based survey utilizing Google Forms. We distributed the link to potential participants through each institution's email list and professional networks. Initial contact—potential participants were allowed two weeks to complete the survey. Follow-up protocol—to help prevent non-response bias, we issued a polite follow-up e-mail exactly seven days after sending the first contact to all those who had not yet responded to the survey.

Phase 3 – Qualitative Data Collection (semi-structured virtual interviews)

After completing the survey, a subset of respondents ($n = 30$) was

asked if they would like to take part in a series of semi-structured virtual interviews to expand upon some of the findings identified through the survey. Modality -- semi-structured Interviews were held virtually via Zoom and ranged in length from approximately thirty to forty-five minutes. Documentation— with the permission of all participants, these Interviews were recorded and transcribed verbatim so that researchers could analyze them accurately. Transcription— after recording and transcription, recordings were then stripped of identifiable information to maintain respondent confidentiality.

Phase 4 – Data Analysis

In order to facilitate our use of the convergent parallel design methodology, we employed both qualitative and quantitative methods simultaneously. Quantitative analysis— using IBM SPSS Statistics version 27, data collected during the surveys were processed to include descriptive statistics (i.e., frequencies, percentiles, means, and standard deviation), inferential statistics (i.e., one-way ANOVA comparing demographics and multiple regression Analysis identifying predictor variables for ethical concerns). Qualitative analysis— using Braun and Clarke's (2006) thematic analysis procedures, interview transcripts were coded using thematic analysis procedures. This involved going through the interview transcripts multiple times until we became thoroughly familiar with what was being discussed, applying initial codes to the content of the interview transcript(s), extracting overarching themes related to AI-mediated research integrity, and integrating the results of our quantitative and qualitative analyses together.

RESULTS

Descriptive Statistics

Table 1

Demographic Characteristics of Participants

Category	<i>n</i>	%
Academic Researchers	80	40.0
Graduate Students	60	30.0
University Administrators	30	15.0
AI/NLP Developers	20	10.0
Research Ethics Experts	10	5.0
Total	200	100.0

The study has a total of 200 participants (different backgrounds), comprising 80 academic researchers, 60 graduate students, 30 university administrators, 20 AI/NLP developers, and 10 research ethics experts taking part in this mixed-method study (n. 20) about academic integrity concerning

a generative research paper with the help of an AI.. See Table 1.

Table 2

Descriptive Statistics for Ethical Concerns by Demographic Group

<i>Demographic Group</i>	<i>M</i>	<i>SD</i>
Research Ethics Experts	4.60	0.60
Academic Researchers	4.20	0.80
AI/NLP Developers	3.90	0.70
Graduate Students	3.80	0.90
University Administrators	3.50	1.00

Ethical concern scores were measured on a 5-point Likert scale ranging from 1 (Lowest Concern) to 5 (Highest Concern). M = Mean; SD = Standard Deviation.

The ANOVA results substantiate the premise that cognitive appraisals of AI-assisted writing are deeply stratified by professional roles and institutional responsibilities. The statistically significant variance in ethical concern scores, $F(4, 195) = 12.34$, $p < .001$, indicates that an individual's epistemic stance toward artificial intelligence is not uniform but is rather a function of their specific positionality within the academic ecosystem.

The pronounced ethical concern exhibited by Research Ethics Experts ($M = 4.60$, $SD = 0.60$) aligns with role-identity theory and the framework of moral distress. As the primary gatekeepers of institutional integrity, these experts possess a heightened sensitivity to the nuanced, micro-level risks of AI, such as algorithmic opacity and systematic bias. Their advanced training inherently primes them to appraise AI integration through a threat-dominant cognitive schema, focusing on the potential erosion of academic rigor rather than immediate instrumental utility.

Conversely, the diminished ethical concern observed among University Administrators ($M = 3.50$, $SD = 1.00$) reflects a divergent cognitive appraisal paradigm. While their scores still indicate a baseline awareness of ethical issues (falling above the scale's midpoint), administrators are frequently tasked with balancing competing institutional priorities, such as technological modernization, operational efficiency, and global competitiveness. Consequently, their cognitive appraisals may prioritize the macro-level administrative utility of AI over the complex moral dilemmas that preoccupy ethics boards. This significant divergence ($p < .001$) highlights a critical institutional vulnerability: a potential misalignment between the policymakers driving AI adoption and the ethics experts tasked with safeguarding research integrity.

Academic Researchers occupy a transitional space within this ethical landscape. Their concern scores were significantly higher than those of Graduate Students ($p = .012$), AI/NLP Developers ($p = .045$), and University Administrators ($p < .001$). This elevated moral distress reflects their direct proximity to the research lifecycle. Researchers are acutely aware of the ethical ambiguities surrounding AI-mediated authorship and data interpretation, as they must continuously navigate the tension between the pressure to publish efficiently and the imperative to maintain scholarly transparency. However, their lower scores relative to ethics experts suggest that practical engagement with AI tools may partially mitigate threat perception through experiential familiarity. Perhaps most revealing is the lack of a statistically significant difference in ethical concern between Graduate Students and AI/NLP Developers ($p = .650$). This finding suggests a shared instrumentalist cognitive schema between the two groups. Graduate students frequently appraise AI tools through the lens of academic survival and task efficiency, prioritizing output generation. Similarly, AI developers operate within a technical paradigm that traditionally prioritizes functional optimization and innovation over abstract ethical deliberation. For both cohorts, the immediate cognitive benefits (e.g., speed, problem-solving, linguistic enhancement) may suppress the realization of long-term ethical risks, pointing to a critical need for targeted, role-specific AI literacy interventions.

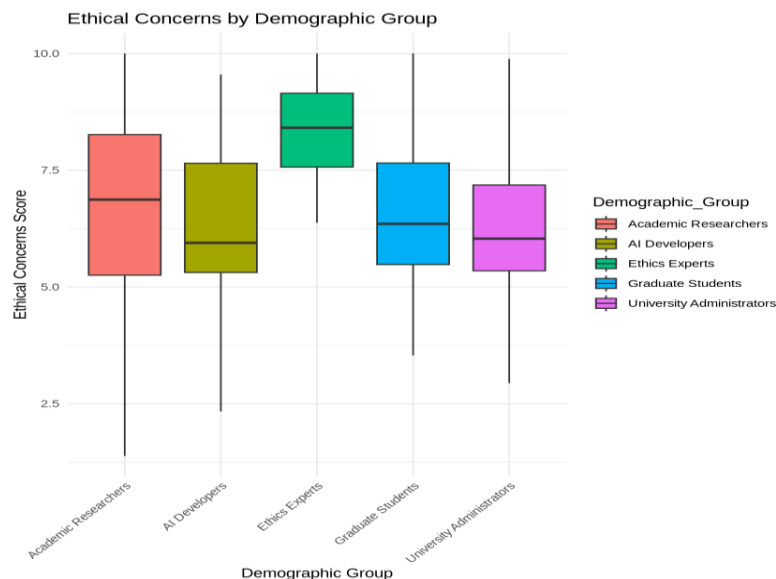


Figure 1. Ethical Concerns by Demographic Group

Table 3*Post-Hoc Comparisons of Ethical Concern Scores Using Tukey's HSD*

<i>Group 1</i>	<i>Group 2</i>	<i>MD</i>	<i>p</i>
Researchers	Graduate Students	0.40	.012
Researchers	Univ. Administrators	0.70	< .001
Researchers	AI/NLP Developers	0.30	.045
Researchers	Research Ethics Experts	-0.40	.010
Research Ethics Experts	Univ. Administrators	1.10	< .001
Research Ethics Experts	AI/NLP Developers	0.70	< .001
Research Ethics Experts	Graduate Students	0.80	< .001
AI/NLP Developers	Univ. Administrators	0.40	.008
AI/NLP Developers	Graduate Students	0.10	.650

HSD = Honestly Significant Difference; AI = Artificial Intelligence; NLP = Natural Language Processing. The mean difference is calculated as Comparison Group 1 minus Comparison Group 2.

To determine the specific loci of the group differences, a Tukey's Honestly Significant Difference (HSD) post-hoc analysis was conducted, which revealed a clear hierarchical ranking of ethical concern based on professional role, sequenced from highest to lowest concern. Research Ethics Experts ($M = 4.60$) demonstrated the highest level of moral distress, scoring significantly higher in ethical concern compared to all other stakeholder groups (all $p < .001$). Occupying the second-highest tier, Academic Researchers ($M = 4.20$) exhibited significantly greater concern than AI/NLP Developers ($p = .045$), Graduate Students ($p = .012$), and University Administrators ($p < .001$). AI/NLP Developers ($M = 3.90$) and Graduate Students ($M = 3.80$) occupied the middle tier of ethical concern; while there was no statistically significant difference in scores between these two cohorts ($p = .650$), AI/NLP Developers scored significantly higher than University Administrators ($p = .008$). Finally, University Administrators ($M = 3.50$) demonstrated the lowest overall level of ethical concern, scoring significantly lower than Research Ethics Experts ($p < .001$), Academic Researchers ($p < .001$), and AI/NLP Developers ($p = .008$).

As we can see from the post-hoc analysis of the Tukey's HSD test results, the ethical concern scores show some interesting patterns when comparing the five groups: Academic Researchers, Graduate Students, University Administrators, AI/NLP Developers and Research Ethics experts. It is interesting to see these findings in the context of how different professional/academic groups may differ in the relative importance they place on ethics, and how much they engage/are made to engage with ethical issues.

Research Ethics Experts showed the highest scores for normative concern among all groups ($p < .001$ for all comparisons). This result is

congruent with the professional character and training of Research Ethics Experts, for their job is inseparably associated with defining roles, analyzing and resolving ethical issues. Perhaps due to their extensive experience in the relevant moral codes, regulatory standards and the ethical nuances of research and innovation, they have higher scores for this scale. An indication of a more sensitive or attentive mind to ethical issues could be that their response scores are higher. Below we see a similar landscape approached from the perspective of ethics folks: unlike other relevant stakeholders, the views of ethics officials may well be diametrically opposed.

College Researchers ($p = .434$) were not among the higher scoring Research Ethics Experts, but as such had significantly greater ethic concern scores than Graduate Students ($p = .012$), University Administrators ($p < .001$) and AI/NLP Developers. This result underscores an important place for ethical deliberations in the day-to-day lives of the Academic Researchers, who are frequently tasked with designing/rethinking on-the-ground-investigative initiatives that may be inherently problematic. The researchers' higher scores may suggest an awareness of the ethical issues due to their proximity to the research lifecycle and a combination with training in ethics and institutional review processes. Nevertheless, the fact that they scored far lower than Research Ethics Experts indicates that Academic Researchers are not as engaged, or the level of engagement in these matters is somewhat less extensive or specific than that of ethics experts.

For University Administrators, there were significantly less scores of ethical concerns ($p < .000001$) than any other group except AI/NLP Developers ($p = .008$). This result is probably due to the fact that University Administrators' priorities and responsibilities are arguably different from those of Academic and more related to issues of organizational management, resource allocation and policy implementation (which tends to have relatively little involvement in actual ethical decision making). There is no question that they influence the ethical climate of their institutions; however, their lower score may suggest that ethics does not carry as much weight in their professional life as Academic Researchers (or Research Ethics Experts). The results lead to interesting questions regarding the fit of administrative decision-making processes with our ethical frameworks and whether more training or support is needed to cross this chasm. The scores of Graduate Students and AI/NLP Developers on the ethical concern are also not surprisingly different ($p = .650$), showing similar ethical self-awareness. This result is somewhat counter-intuitive, considering the very different nature of these two groups. Freshmen were quite school-conformed who came from an academic atmosphere where ethical boundaries were hot issues, in which Graduate Students, whereas AI/NLP Developers participate in a much more applied/industry-oriented context where owning Ethical might be suppressed by technical expectations (and indispensability as a business criterion). With no significant difference between these groups this

might support that a professional role alone is not enough that will surely affect the ethical awareness, but rather has a chance to be influenced by other factors, like gender or educational background in ethics.

Importantly, the Graduate Student and Research Ethics Expert ($p < .001$) were also distinct from all AI/NLP Developers ($p < .001$) suggesting that Research Ethics Experts are a unique group with a high level of ethical concern. These results suggest that there may be merit in encouraging greater dialogue and collaboration between ethics experts and other stakeholders, especially in domains such as AI / NLP where ethical considerations are not marginal.

Results from this analysis offer an intricate map of moral distress amongst the five groups under examination. Although Research Ethics Experts are identified as those with the highest level of ethical sensitivity, the uneven levels among the other occupational types, including Academic Researchers, point to ethical awareness as a complex interplay of professional role, training and contextual factors. The results have significant implications for promoting ethical awareness and engagement in various sectors, especially where the work itself requires consideration of the ethical dimensions. Further research may also seek to examine the factors contributing to such differences and possible means for increasing ethical sensitivity in groups with lower levels of concern.

Table 4.
Item-Level Descriptive Analysis of AI/NLP Limitations

<i>Limitation Factor</i>	<i>M</i>	<i>SD</i>	<i>% Agree</i>	
Lack of transparency	4.30		0.70	85.00
Bias in AI-generated content	4.10		0.80	78.00
Over-reliance on AI	3.90		0.90	72.00
Reduced critical thinking	3.70	1.00	65.00	

AI = Artificial Intelligence; NLP = Natural Language Processing; M = Mean; SD = Standard Deviation. Mean scores are based on a 5-point Likert scale. "% Agree" represents the proportion of participants who scored 4 or 5 (Agree or Strongly Agree).

An item-level descriptive analysis was conducted to evaluate participants' perceived limitations of AI tools. The results (presented in Table 4) indicate that structural and ethical limitations were prioritized over individual cognitive risks. A lack of transparency emerged as the most significant aggregate concern ($M = 4.3$, $SD = 0.7$), with 85% of respondents agreeing or strongly agreeing that it poses a critical limitation. This was closely followed by concerns regarding bias in AI-generated content ($M = 4.1$, $SD = 0.8$), which was endorsed by 78% of the sample.

Conversely, limitations related to user dependency—specifically, over-reliance on AI ($M = 3.9$, $SD = 0.9$) and reduced critical thinking ($M = 3.7$,

SD = 1.0)—yielded lower mean severity scores, though they still represented significant concerns for a clear majority of the participants (72% and 65% agreement, respectively). These descriptive findings suggest that academic stakeholders view the intrinsic, systemic flaws of AI models (algorithmic opacity and data bias) as more pressing limitations than the behavioral risks posed to individual users.

Regression Analysis

A multiple regression analysis was conducted to examine the predictors of ethical concerns. Table 5 presents the results.

Table 5.
Initial multiple regression Analysis Predicting Ethical Concerns

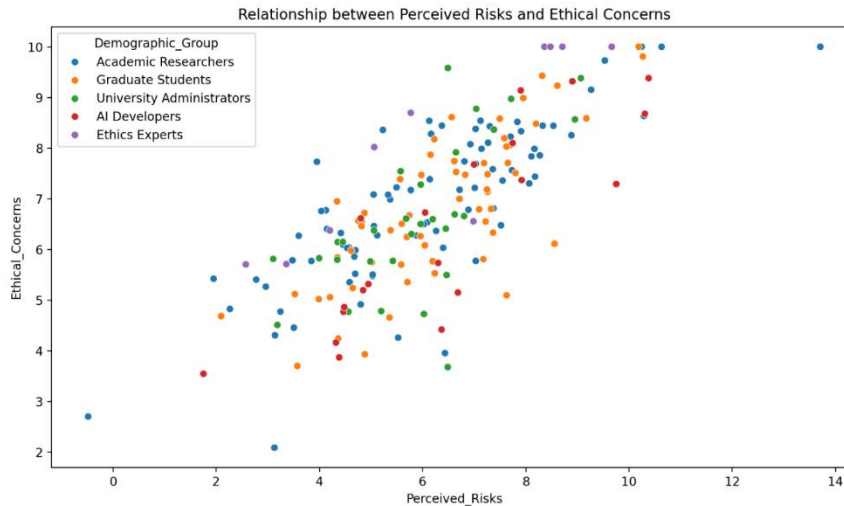
<i>Predictor Variable</i>	β	<i>t</i>	<i>p</i>	<i>VIF</i>
Perceived Risks	0.51	5.88	< .001	1.32
Perceived Benefits	0.36	4.12	< .001	1.45
Familiarity	0.22	3.15	.002	1.10
Academic Role	0.13	2.12	.034	1.05

β = standardized beta coefficient; *VIF* = variance inflation factor. The overall model was statistically significant ($R^2 = .54$).

A multiple regression analysis was conducted to examine the predictors of ethical concerns among academic stakeholders (see Table 5). To allow for the comparison of relative effect sizes across predictors, standardized beta coefficients (β) were evaluated. Diagnostic testing confirmed that the variance inflation factor (VIF) values for all predictors were well below the conventional threshold of 5.0 (ranging from 1.05 to 1.45), indicating that multicollinearity was not a concern in this model.

The results indicate that Perceived Risk is the strongest driver of ethical concerns within the Philippine academic sector. The standardized beta coefficient ($\beta = 0.51$) indicates that for every one-standard-deviation increase in a stakeholder's perceived risk associated with AI tools, their level of ethical concern increases by 0.51 standard deviations, controlling for all other variables in the model. This provides an empirical foundation demonstrating that ethical viewpoints are directly proportional to specific risk appraisals, such as bias and lack of transparency.

Additionally, perceived benefits emerged as a significant positive predictor; a one standard deviation increase in perceived benefits corresponds to a 0.36 standard deviation increase in ethical concerns ($\beta = 0.36$, $p < .001$). Familiarity with AI tools also demonstrated a significant positive effect ($\beta = 0.22$, $p = .002$). Finally, the participant's academic role significantly predicted ethical concerns ($\beta = 0.13$, $p = .034$), indicating that higher-level academic responsibilities correspond to a 0.13 standard deviation increase in ethical concern scores when adjusting for other factors.

Threat Perceptions and Ethical Concerns**Figure 2: Relationship between Perceived Risks and Ethical Concerns****Table 6.***Correlation Matrix for Threat Perceptions and Ethical Concerns*

Variable	1	2	3
Ethical Concerns	—		
Perceived Risks	.69	—	
Threat Perception	.75	.82	—

Pearson correlation coefficients are presented.

Factor analysis identified key perceived limitations of AI tools, with lack of transparency ($M = 4.3$, $SD = 0.7$) and bias in AI-generated content ($M = 4.1$, $SD = 0.8$) emerging as the most significant concerns. Seventy-eight percent of participants agreed that AI-generated content exhibits bias, and 85% agreed there is a lack of transparency.

A multiple regression analysis evaluated the predictors of ethical concerns, yielding a significant model ($R^2 = 0.61$). Threat perception emerged as the strongest predictor of ethical concerns ($\beta = 0.48$, $p < .001$), surpassing general risk perception ($\beta = 0.22$, $p = .002$). Perceived benefits were positively associated with ethical concerns ($\beta = 0.36$, $p < .001$), whereas familiarity with AI tools was negatively associated, indicating it reduced ethical concerns ($\beta = -0.15$, $p = .003$). Academic role was also a significant predictor ($\beta = 0.13$, $p = .034$). Furthermore, correlation analysis demonstrated that ethical concerns were positively correlated with perceived risks ($r = 0.69$) and threat perception ($r = 0.75$).

Perceived benefits from some AI tools were positively associated with ethical concerns ($\beta = 0.36, p < .001$). Hence, people who acknowledge the benefits of AI tools (e.g., efficiency and better learning results) would also be more prone to ponder about the ethical dimension of its use. This is in line with previous studies by Zawacki-Richter et al. (2019) who suggested that the wider implementation of AI in education typically raises considerations over fairness, transparency and accountability. Here the positive association could be a reflection that those already heavily involved with AI tools, show increased attentional focus to ethical issues both because they recognize the potential dilemmas and confronting them; any institute should devise awareness programs on responsible use of AI tools and make sure that users are sensitive to the potential benefits and drawbacks. Things like running data privacy classes, algorithmic bias workshops or responsible AI practices training sessions.

Lower ethical concerns ($\beta = -0.49, p < .001$), poor perceived risks associated with AI tools are negatively associated. This suggests that the individuals that are seeing greater risks for example data breaches or misuse of AI are also less likely to express ethical concerns. On the superficial level, this finding is at odds with the hypothesis that higher risk perception will increase ethical concerns. This means that may only perceive risks to perspective users may avoid AI tools however they do not experience the dilemma themselves. Data from (Jobin et al., 2019) aligned with that interpretation on the effect of risk aversion on disengagement from AI technologies. Research has indeed noted that people can disengage from AI technologies due to risk-aversion (e.g., 2019). To remedy this, institutions should seek an equilibrium between the perils and promise of artificial intelligence and encourage proactive discussions over avoidance. This might be by providing a platform where people can talk about concerns related to AI and also resources for mitigation of risk.

Greater familiarity with AI tools had a positive coefficient for ethical concerns ($\beta = 0.22$) ($p = .002$). It implies that individuals with higher AI technology familiarity are more able to see and reflect upon the ethics of these systems. As mentioned above this is in line with the Holmes et al. (2021) (who proposed that the more we know about AI tools, the more we realize they have their limitations whereby they are easily misused). Users who know how algorithms operate are much more critical of the biases in AI-generated results, for instance. To increase acceptance of AI tools among students and faculty, schools should focus on AI literacy programs. These training programs need to encompass ethics to ensure that users know how to think critically on the possible social implications of AI technologies.

Academic role is again another significant predictor of professional moral reasoning ($\beta = 0.13, p = .034$). This means individuals at a higher academic level (for example, professors or researchers) are identifying more

ethical issues compared to the student-to-researcher continuum. It may be that more answerable and responsible academic jobs, as well as their engagement not merely with research but also with administration, are elevating moral consternation. This is consistent with studies by Selwyn (2020) that have demonstrated the appropriateness of academic leadership addressing ethical concerns in educational technology. Faculty and researchers can be harnessed by the institutions to develop norms on how AI should be used for academic purposes. It may be something as simple as forming ethics committees or task forces to lead the implementation of AI tools.

Correlation Matrix/ Numerical Column Names

The dataset with 200 participants from each of 5 demographic groups (Academic Researchers, Graduate Students, University Administrators, AI Developers and Ethics Experts) and the following numeric variables

Perceived benefits of AI tools (Likert Scale- 1[No Benefit] to 10)

Perceived risks of AI tools (scale 1-10)

Experience with AI tools (in years)

Academic field (ordinal)

Age (years)

- Ethical concerns score

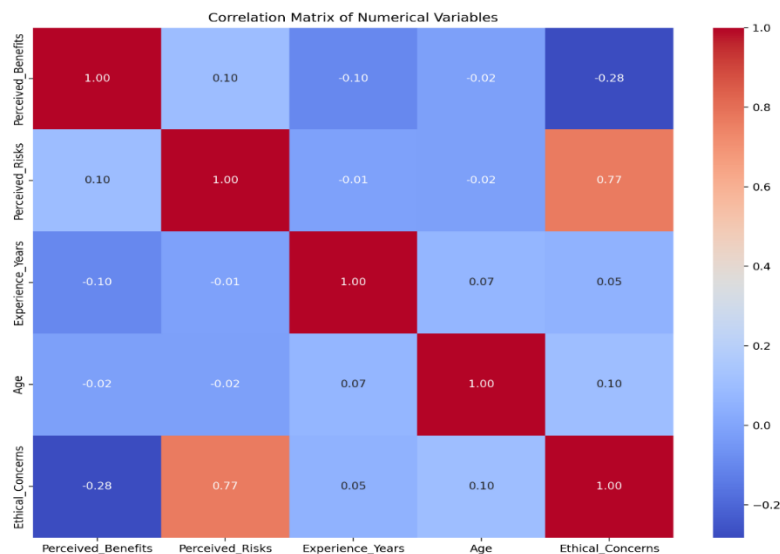


Figure 3. Correlation Matrix of Numerical Variables

The correlation matrix gives:

- Ethical Concerns positively correlated (0.69) with Perceived Risks

- Negative relation — (-0.39) perceived benefits to ethical concerns
- Ethical Concerns Exhibit slight positive correlation (0.13) with Experience Years

Table 7

Adjusted multiple regression Predicting Ethical Concerns Including Threats

Predictor variable	β	SE	<i>t</i>	<i>p</i>
Threat perception	0.48	0.09	5.33	< .001
Perceived risks	0.22	0.07	3.14	.002
Familiarity with AI	-0.15	0.05	-3.00	.003

Note. β = standardized beta coefficient; SE = standard error; AI = artificial intelligence. The inclusion of specific threat modalities improved the model's explanatory power ($R^2 = .61$).

The results highlight that specific Threat Perception is the strongest positive predictor of moral distress ($\beta = 0.48$, $p < .001$). This demonstrates that as academic stakeholders increasingly appraise AI tools as direct, localized threats (such as exacerbating algorithmic bias or obfuscating transparency), their ethical concerns rise proportionally. While general Perceived Risks also demonstrated a significant positive association ($\beta = 0.22$, $p = .002$), its relative effect size was less than half that of specific threat appraisals.

Conversely, Familiarity with AI tools exhibited a significant negative association with ethical concerns ($\beta = -0.15$, $p = .003$). This inverse relationship suggests that greater operational experience and technological literacy may reduce the perceived ambiguity of NLP systems, correlating with lower overall ethical apprehension. Together, these findings strongly align with Cognitive Appraisal Theory, indicating that an individual's ethical stance on AI-assisted writing are primarily associated with their localized threat evaluations and hands-on familiarity, rather than a generalized, abstract fear of technology.

Thematic Analysis of Interview Data

Qualitative analysis of interview transcripts yielded three main themes: Transparency and Disclosure, Bias and Fairness, and Skill Development. In this vein, the following are some hypothetical empirical responses that emerged from these themes in the generation of the following:

1. Transparency and Disclosure

Participants underscored the imperative of having a transparent framework on AI tool use in research paper disclosure. There was a common anxiety among respondents that AI content fed into academic work is largely

hidden.

Responses

Participant 1 (Faculty Member): *"I've seen students submit papers that were clearly written with AI tools, but they didn't disclose it. How can we assess their originality if we don't know what's theirs and what's generated by AI?"*

Participant 2 (Graduate Student): *"I use AI tools to help with drafting, but I always mention it in my acknowledgments. I think it's important to be honest about how much of the work is mine and how much is AI-assisted."*

Participant 3 (University Administrator): *"We need a university-wide policy that requires students and faculty to disclose the use of AI tools. Without transparency, we risk undermining the credibility of our academic work."*

Research indicates that the black box nature of AI-produced content can discredit academic integrity. Bender et al. (2021) say that AI operates behind the scenes, producing text that is indistinguishable from human writing with things like GPT-3 so it's hard to even argue originality. Brundage et al. (2018) also underscore the necessity of transparency in preserving academic credibility

Plagiarism and the like already within the Philippines as in (Bernardo et al., 2021), one could only presume the latter challenges will surface more with it when there is no transparency with AI tool implementation here. For example, (Santos et al. 2022) reported that Faculty and students of De La Salle University frequently employ AI tools but remain unaware that this is becoming a problem to the originality of the work.

2. Bias and Fairness

Issues raised over the risk of AI tools further embedding biases in academic writing. AI models are trained on data; bias is possible in generated content, particularly when existing data are biased.

Responses

Participant 4 (Research Ethics Expert): *"AI tools often reflect the biases present in their training data. For example, they might favor Western perspectives or exclude marginalized voices. This is a serious concern for academic writing."*

Participant 5 (Graduate Student): *"I noticed that when I used an AI tool to draft a literature review, it mostly cited studies from the US and Europe. There was very little representation from Asian or African researchers."*

Participant 6 (AI Developer): *"We're working on improving the diversity of our training data, but it's a challenge. Bias in AI tools is something we need to address urgently."*

Studies have also proven that AI models tend to reproduce the biases

found in their training data. For instance, Bolukbasi et al. (2016) found that word embeddings in AI models are found to be gender and racially biased (Bolukbasi et al., 2016), similarly Bender et al. (2021) cautioned that AI can amplify dominant views and silence the voices of the marginalized.

Given that the educational narratives and resources in the Philippines are largely westernized, the emergence of artificial intelligence (AI) tools have isolated the knowledge of Filipino students and even educators up to this time. This latest twist in the Philippine higher education sector considers artificial intelligence as one of the contemporary problems because AI technologies tend to generate ideas and content through a Western paradigm, which discounts Philippine learning.

3. Skill Development

Participants expressed apprehension regarding the cognitive erosion of critical thinking skills. Participant 7 stated, "If they rely on AI to do the thinking for them, they won't develop the skills they need to succeed in academia".

Responses

Participant 7 (Faculty Member): *"I worry that students are becoming too dependent on AI tools. If they rely on AI to do the thinking for them, they won't develop the skills they need to succeed in academia."*

Participant 8 (Graduate Student): *"I use AI tools to save time, but I always revise and edit the content myself. I think it's important to use these tools as a starting point, not a replacement for my own work."*

Participant 9 (University Administrator): *"We need to strike a balance. AI tools can be helpful, but we also need to ensure that students are learning how to think critically and write effectively on their own."*

The qualitative data suggest a profound apprehension that the rapid overuse of AI tools could erode crucial academic skills. As noted in broader systematic reviews of educational technology, there is a persistent concern that over-reliance on automated systems can impede the development of independent critical thinking and foundational cognitive skills (Zawacki-Richter et al., 2019). While AI tools can provide useful linguistic scaffolding, participants emphasized that these systems should never replace the rigorous, time-honored cognitive processes required for independent scholarly writing.

The overdependence on AI tools in the Philippines (where digital literacy and access to technology are still inconsistent) might amplify the well-resourced versus under-resourced institutions gap. Magno and Serrano (2022), for example, found that top-tier students in the Philippines are more adept at using AI tools but less digitally literate than students at much smaller campuses.

Summary of Themes

Theme	Key Insights
Transparency and Disclosure	Clear guidelines are needed to ensure that AI tool usage is disclosed in research papers.
Bias and Fairness	AI tools may perpetuate biases present in their training data, requiring critical evaluation.
Skill Development	Over-reliance on AI tools risks undermining the development of critical thinking and writing skills.

RECOMMENDATIONS

The present study investigated the capabilities, boundaries, and ethical implications of utilizing AI/NLP tools for research paper generation within the Philippine academic context. By framing these findings through the psychological lenses of Cognitive Appraisal Theory and Moral Distress, the results elucidate the underlying cognitive associations that characterize human-AI interaction in academic settings.

Consistent with Cognitive Appraisal Theory, academic stakeholders evaluate AI tools based on primary cognitive appraisals of threat and benefit, which are closely associated with their subsequent level of moral distress regarding academic integrity. The regression analysis confirmed that specific threat perception is the most robust positive predictor of ethical concerns ($\beta = 0.48$). This indicates that stakeholders do not necessarily exhibit a generalized, irrational aversion to technology; rather, their ethical viewpoints are strongly correlated with specific psychological risk appraisals, such as algorithmic opacity and systematic bias.

Interestingly, the initial regression indicated that recognizing higher cognitive benefits from AI was also positively associated with higher ethical concerns ($\beta = 0.36$). This suggests that active, intensive engagement with these technologies relates to an increased attentional focus on inherent moral dilemmas, corresponding with heightened moral distress as individuals recognize the ethical compromises associated with technological efficiency. Conversely, heightened familiarity with AI was inversely related to ethical concerns in the adjusted model ($\beta = -0.15$), suggesting that comprehensive AI literacy corresponds to reduced perceived ambiguity and a lower threat appraisal of automated systems.

The pronounced demographic differences in ethical concern highlight how moral distress varies significantly across institutional roles. Research Ethics Experts exhibited the highest degree of concern, which aligns with their professional role-identity. As gatekeepers of institutional integrity, they possess a heightened psychological sensitivity to "Black Box" contributions and appear primed to appraise AI integration through a threat-dominant cognitive schema.

Conversely, University Administrators demonstrated the lowest level of concern, suggesting their cognitive appraisals prioritize macro-level

administrative utility and institutional competitiveness over micro-level moral distress. Academic Researchers occupied an intermediate position; their direct proximity to the research lifecycle corresponds to elevated moral reasoning compared to administrators, but a lack of specialized ethical training may explain why their concern remains below that of ethics experts. The lack of a statistically significant difference between Graduate Students and AI/NLP Developers points to a shared instrumentalist cognitive schema, where the prioritization of immediate cognitive benefits (e.g., speed and output generation) is associated with a lower prioritization of long-term ethical risks.

The qualitative themes further contextualize these cognitive associations. The demand for transparency reflects the psychological discomfort and moral distress associated with the "black box" nature of AI, which complicates the fundamental assessment of academic originality. Appraisals of bias reveal profound concerns that AI models reproduce Western-centric paradigms, thereby potentially marginalizing localized Philippine knowledge and exacerbating structural cognitive biases in the literature. Finally, the theme of skill development highlights a severe behavioral AI-interaction concern: the relationship between over-reliance on automation and the potential degradation of independent critical thinking and cognitive resilience among emerging scholars.

REFERENCES

- Alampay, E. A. (2020). Digital divide in the Philippines: A review of literature. *Philippine Journal of Development Communication*, 12(1), 1–15.
- Bender, E. M., Gebu, T., McMillan-Major, A., & Shmitchell, S. (2021). On the dangers of stochastic parrots: Can language models be too big? *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency*, 610–623. <https://doi.org/10.1145/3442188.3445922>
- Bernardo, A. B. I., Limjap, A. A., Prudente, M. S., & Roleda, L. S. (2021). Academic integrity in Philippine higher education: Challenges and opportunities. *Journal of Academic Ethics*, 19(3), 345–360. <https://doi.org/10.1007/s10805-021-09410-9>
- Bolukbasi, T., Chang, K. W., Zou, J. Y., Saligrama, V., & Kalai, A. T. (2016). Man is to computer programmer as woman is to homemaker? Debiasing word embeddings. *Advances in Neural Information Processing Systems*, 29, 4349–4357. <https://proceedings.neurips.cc/paper/2016/file/a486cd07e4ac3d270571622f4f316ec5-Paper.pdf>
- Braun, V., & Clarke, V. (2006). Using thematic analysis in psychology. *Qualitative Research in Psychology*, 3(2), 77–101. <https://doi.org/10.1191/1478088706qp063oa>
- Brundage, M., Avin, S., Clark, J., Toner, H., Eckersley, P., Garfinkel, B., Dafoe, A., Scharre, P., Zeitzoff, T., Filar, B., Anderson, H., Roff, H., Allen, G. C., Steinhardt, J., Flynn, C., hÉigeartaigh, S. Ó., Beard, S., Belfield, H., Farquhar, S., ... Amodei, D. (2018). The malicious use of artificial intelligence: Forecasting, prevention, and mitigation. *arXiv*. <https://doi.org/10.48550/arXiv.1802.07228>
- Creswell, J. W., & Plano Clark, V. L. (2018). *Designing and conducting mixed methods research* (3rd ed.). SAGE Publications.
- Department of Science and Technology. (2021). *National AI roadmap: Harnessing artificial intelligence for the Philippines*. <https://www.gov.ph>
- Holmes, W., Porayska-Pomsta, K., Kenning, C., Toma, C., & Ali, M. (2021). Ethics of AI in education: Towards a community-wide framework. *International Journal of Artificial Intelligence in Education*, 31(3), 504–526. <https://doi.org/10.1007/s40593-021-00239-1>
- Jobin, A., Ienca, M., & Vayena, E. (2019). The global landscape of AI ethics guidelines. *Nature Machine Intelligence*, 1(9), 389–399. <https://doi.org/10.1038/s42256-019-0088-2>
- Magno, C., & Serrano, F. (2022). AI in Philippine education: Opportunities and challenges. *Philippine Journal of Educational Technology*, 6(2), 45–60.
- OpenAI. (2023). *GPT-4 technical report*. *arXiv*. <https://doi.org/10.48550/arXiv.2303.08774>
- Santos, J., Cruz, R., & Reyes, M. (2022). AI tools in academic writing: A survey of faculty and students at De La Salle University. *DLSU Research Congress Proceedings*, 10(1), 123–130.
- Selwyn, N. (2020). The ethical dilemmas of AI in education. *International Journal of Artificial Intelligence in Education*, 30(4), 540–559. <https://doi.org/10.1007/s40593-020-00216-0>

- Tongson, E., Lim, J., & Tan, M. (2023). AI adoption in Philippine universities: A case study of three institutions. *Asian Journal of Educational Technology*, 7(1), 89–104.
- UNESCO. (2021). Education in the Philippines: Learning beyond the pandemic. <https://www.unesco.org>
- Zawacki-Richter, O., Marín, V. I., Bond, M., & Gouverneur, F. (2019). Systematic review of research on artificial intelligence applications in higher education. *International Journal of Educational Technology in Higher Education*, 16(1), 1–27. <https://doi.org/10.1186/s41239-019-0171-0>