

Exploring User Creativity and Ethical Decision Making in Generative AI: A Transformer-Based Semantic and Network Analysis of Reddit Discussions on ChatGPT

MSc Research Project
Master of Science in Data Analytics Information

Gaurav Sadanand Satam
Student ID: 23269898

School of Computing
National College of Ireland

Supervisor: Cristina Hava Muntean

National College of Ireland
MSc Project Submission Sheet
School of Computing



Student Name: Gaurav Sadanand Satam
Student ID: 23269898
Programme: M.Sc. Data Analytics **Year:** 2025-2026
Module: Research Practicum
Supervisor: Cristina Hava Muntean
Submission Due Date: 28-01-2026
Project Title: Exploring User Creativity and Ethical Decision Making in Generative AI: A Transformer-Based Semantic and Network Analysis of Reddit Discussions on ChatGPT
Word Count: **21 Pages** **Count 9435 words**

I hereby certify that the information contained in this (my submission) is information pertaining to research I conducted for this project. All information other than my own contribution will be fully referenced and listed in the relevant bibliography section at the rear of the project.

ALL internet material must be referenced in the bibliography section. Students are required to use the Referencing Standard specified in the report template. To use other author's written or electronic work is illegal (plagiarism) and may result in disciplinary action.

Signature: Gaurav Sadanand Satam

Date: 28-01-2026

PLEASE READ THE FOLLOWING INSTRUCTIONS AND CHECKLIST

Attach a completed copy of this sheet to each project (including multiple copies)	☐
Attach a Moodle submission receipt of the online project submission, to each project (including multiple copies).	☐
You must ensure that you retain a HARD COPY of the project, both for your own reference and in case a project is lost or mislaid. It is not sufficient to keep a copy on computer.	☐

Assignments that are submitted to the Programme Coordinator Office must be placed into the assignment box located outside the office.

Office Use Only	
Signature:	
Date:	
Penalty Applied (if applicable):	

Exploring User Creativity and Ethical Decision Making in Generative AI: A Transformer-Based Semantic and Network Analysis of Reddit Discussions on ChatGPT

Gaurav Sadanand Satam
23269898

Abstract

The creation of generative AI applications, such as ChatGPT, at a very fast pace has changed the balance of interaction between individuals and has already begun to redefine creativity and ethical choice in daily life. The research questions the use of ChatGPT in teaching Reddit communities by students and the inventiveness and ethical behaviour dynamic. The study measures the emotional, ethical, and creative profile of 177,485 posts on Reddit with the help of the How Students Use ChatGPT Filtered Reddit Dataset and an advanced transformer-based pipeline, which combines RoBERTa sentiment analysis, BART-based zero-shot ethical stance classification, and Sentence-BERT semantic embeddings. It is revealed that the student discourse is mostly neutral to positive (neutral 50.5%, positive 31.7%, negative 17.8%), and the majority of the conversations involve the description of pragmatic and informational ChatGPT uses, instead of polarised responses. The prosocial posts (53.3%), and the neutral ones (23.4%), and the posts showing potentially harmful or risky usage (12.8%), are identified as the ethical stance modelling demonstrates that more than half of the posts are prosocial (53.3%), and another significant proportion (23.4 percent) is the neutral ones. Creativity scores show that discussions are more informative than radically innovative with zero-shot labels being dominated by the informative posts (78.1%) and very few posts being highly creative (1.6%), which may suggest that students use ChatGPT more as an intellectual assistant than as the source of original invention. Combining these composite scores with a lightweight co-occurrence network, which identifies central tags like Conversation and College as central discourse nodes, the study has provided a multidimensional analytical paradigm that will contribute to the existing knowledge regarding the impact that generative AI is having on student creativity, ethics, and community norms in the real-world online context.

1 Introduction

In terms of discussing about the background of this research conducted, Generative Artificial Intelligence (GenAI) has been having a significant impact on digital interaction, creative production, and problem-solving practices in many fields in recent years. OpenAI ChatGPT is one of these innovative technologies as a state-of-the-art conversational agent that can read and write texts as a human being. Having the potential to support a wide range of tasks, including content creation, educational support, and more, ChatGPT and other models of this sort have been rapidly adopted, especially among student communities. The rapid adoption of GenAI by students has made much of the world interested in how these tools affect creativity, decision-making and ethical consciousness in academic and real-life situations. With the help of new tools of GenAI, such as ChatGPT, users can familiarize themselves with new ideas, solve complex tasks and develop various creative works with more ease than ever before. However, despite these good points, some questions about the reliance of users to AI-generated content, the potential disappearance of original thinking and ethical considerations of the application of such intelligent systems are yet to be answered. This rapid progress of GenAI demands severe and evidence-based research of its impact, especially how individuals orient on creativity and moral that are connected to AI assistance in the reality of everyday events.

The research problem provides information on the main issue that this study aims to resolve is to explain the subtle effects of generative AI applications on the creativity and moral judgment abilities of users, especially when young learners use ChatGPT in social media sites like Reddit. Although user sentiment and topic trends related to ChatGPT have been studied previously and discussed in

detailed in the ‘Literature Review’ section of this report, it is stated that there is a lack of the research that quantifies creativity and ethical considerations that users reflect on in a systematic manner. Furthermore, most of the current analyses are usually based on conventional methods such as sentiment analysis and Latent Dirichlet Allocation (LDA) topic modelling that, though informative, do not represent the inherent semantic nuances and ethical positions held by the users. This study will fill that gap by using sophisticated transformer-based natural language processing (NLP) methods to extract deeper thematic knowledge, ethical positions, and creativity indicators in the user-generated material. The difficulty is not to only interpret large-scale social media data but also create rigorous metrics by which inventiveness and ethical awareness can objectively be measured as reflected in conversational AI interactions.

Motivation gained to perform this research is due to the increased diffusion of generative AI systems into learning settings highlights the immediate necessity to understand their more global cognitive and ethical consequences. The insights into the application of ChatGPT by students provide an insight into the opportunities and risks of incorporating AI into the learning process, creative problem-solving, and decision-making. The research aims at informing teachers, policy makers, and the creators of AI tools about the practical consequences of GenAI technologies on the reasoning and ethics of youth by dissecting ethical considerations and expressions of creativity in online discourse. Also, the desire to introduce a new methodological paradigm based on the integration of semantic analysis, the identification of ethical positions, the social network mapping, and the mixed qualitative-quantitative research is the driving force of this study. This kind of framework will potentially yield more valuable insights that can be applied to practice than traditional sentiment or topic categorization techniques alone, eventually leading to responsible AI usage and better design of generative tools that are sensitive to human creativity and ethics.

On the basis of the study performed, the research question that this study tends to answer is, “What are the ways to measure creativity, ethical position, and sentiment in Reddit-based discussions about Generative AI tools, such as ChatGPT, through the use of transformer-based models (RoBERTa, Sentence-BERT (MiniLM) and BART) and social network analysis?”

The current study provides a groundbreaking contribution to the literature on Human-AI interaction as it provides well-grounded empirical data on the effect of intensive use of Generative AI-based applications such as ChatGPT on the creativity of the user and their moral decision-making. The study, using state-of-the-art transformer-based natural language processing and social network analysis of massive Reddit discussions, is the first to quantify inventiveness and ethical positions as they occur naturally within digital communities. Not only do the results contribute to the scholarly body of knowledge regarding cognitive augmentation and ethical reasoning in the framework of GenAI, but they also offer practical implications related to the development of AI systems that enhance the transparency of systems and support creativity as well as safeguard ethics in their systems. Besides this, the results are useful to education institutions and policymakers aiming at instilling AI literacy and responsible use of GenAI by designing evidence-based curricula and training in digital ethics. Besides, the project forms a methodological prototype to be used by future researchers by establishing composite indices of creativity and ethics, which indicates how highly developed machine learning can outgrow basic sentiment analysis and provide a more detailed picture of human-AI partnership. Finally, the study enhances our knowledge about creativity, technological reliance, and new sets of norms in collaborative human-AI ecosystems in the academic and societal settings.

2 Structure of Research

The report is structured into a series of sections that proceed to formulate the theoretical framing, methodological design and empirical findings of the study. The Literature Review section discusses the recent research on the topic of generative AI in the creative, educational, and ethical settings, and provides the overview of the existing strategies of studying the large-scale user discourse. The Research Methodology section outlines the general study design, such as the dataset selection, data pre-processing, applying transformer-based NLP models, providing strategies to measure creativity and ethics, and combines the temporal and social network analysis in a mixed-method approach. The Results section contains the key analytical results, including the descriptive summaries, the visualisation of the trends in creativity and ethical tendencies, the network-mapping outputs, and the most important statistical relationships based on the conversations in Reddit. These findings are

interpreted in the Discussion section with reference to the existing Human-AI interaction theory and practice, and new insights, design implications of generative AI tools, and limitations are in critical perspectives. Lastly, the Conclusion section summarises the overall contributions of the work, comments on both methodological and substantive limitations, and gives recommendations on how to use GenAI responsibly and what research should do on the creativity and ethics of AI-mediated environments in the future.

3 Literature Review

Generative artificial intelligence (GenAI) models and chatbots such as ChatGPT have influenced digital communication, creativity and decision-making scenarios in a significant way. With the introduction of GenAI into human life, researchers like Gupta P. et al. and others have been interested in the study of its effects on creativity, ethical judgments, and improvement of cognitive processes (Gupta P. et al., 2024). As per a study conducted by Talafidaryani M. and Moro S., Reddit in particular, is a well-organized source of social media communication data that can be analysed in large volumes with semantics, social, and ethical analysis of user interactions with GenAI (Mojtaba Talafidaryani M. and Moro S., 2024).

3.1 Recent advancement in Semantic Models based on Transformer

In their study, Trigka M. and Dritsas E. explained that with self-attention mechanisms transformers have revolutionized the natural language processing (NLP) process with the ability of a deeper representation of semantics compared to the past neural architectures. Such models as BERT, GPT-3, and GPT-4 have shown best-in-class performance in such tasks as sentiment analysis and text-generation key findings of which have been discussed in this study (Trigka M. and Dritsas E., 2025). Using the example of Sengar S. S. et al. who examined the systems of use of transformer creatively, medically, and educationally and concluded that the systematically adaptable and sentiment-sensitive analyses were highly developed (Sengar S. S. et al., 2024). The success of sentiment analysis on transformer models with both BERT and GPT models in marketing was also established by Petratos P. and Giannoula M., as the classification accuracy improved more than 10 percent in comparison to traditional models, and hence the digital communication analysis in real life is appropriate (Petratos P. and Giannoula M., 2025). It has been demonstrated by Koetz J. et al that transformers have never lagged behind classical NLP models in understanding and generating complex and multidimensional textual content relative to news media and mental health (Koetz J. et al., 2025). However, these studies largely apply traditional standards of accuracy without addressing the fact that it may be affected by biases in training data, which could not only affect interpretability and fairness.

These models are primarily based on big, labelled benchmark datasets and fine-tuning of transformer encoders or encoder-decoder models. They are generally used to compare performance on older models like the LSTMs or TF-IDF + classifiers on problems such as sentiment classification or topic detection. Majority of them report increases of 5-15 percentage point in accuracy / F1, yet they concentrate on predictive performance and do not talk much about explanation or fairness. Not many of these papers use transformers on large, mixed purpose social media corpora with poor labels, as mine does on the Reddit data.

3.2 GenAI, Creativity, and Human-AI Collaboration

One of the themes of GenAI research is creativity. Desdevises J. examined the paradoxes of creativity by performing an experiment involving ChatGPT-4 on the basis of an egg task that measured bias and the originality of ideas presented by the participants. He explained that ChatGPT-4 is very fluent and productive, but it has difficulty differentiating between original and conventional ideas, which could restrict creative thinking (Desdevises J., 2025). Jackson V. et al. measured the effect of GenAI on the creativity of software development empirically and discovered that it enhanced the generation of ideas but with inconsistent originality and adaptability (Jackson V. et al., 2025). According to an analysis of Reddit posts by Buz T. et al., large language models are capable of sounding like humans, but they do not have the subtlety required to be truly creative. These inferences were made after a detailed study conducted by them which concluded that human evaluators often failed to differentiate between human-written and AI generated texts (Buz T. et al., 2024). Chugh R. et al. observed that human-machine teamwork is important in improving creativity because it complements human intelligence as opposed to replacing it (Chugh R. et al., 2025). While on the

other hand, Chiarello F. et al. used a dataset containing 3.8 million tweets during their study in order to highlight the scope of upcoming generative models should serve as a collaborative assistant to enable user creativity by bridging the gap between generative LLMs and tasks assigned to them. Out of the many tweets in their dataset, they discovered that roughly 263,000 tweets were about assigning a task contributing to 6.86% of the total tweets (Chiarello F. et al., 2024). However, these researches usually do not have thorough analyses of the long-term creative impacts and consider only short-term productivity indicators.

The experiments used in the creativity studies typically feature controlled experiments, survey measures or small task-based measures having tens or hundreds of participants. They emphasize short-term idea generation, fluency and originality scores as opposed to large-scale naturally occurring text data. A significant number of them use human raters to grade creativity of outputs subjecting them to scalability and introducing subjectivity. All these works do not calculate semantic or lexical diversity on a systematic basis across hundreds of thousands of posts, which is an idiosyncratic aspect of my method.

3.3 Bias, Ethical Decision-Making and Public Attitudes

GenAI has become an urgent subject of the ethical debate due to its rapid adoption. The ethical issues that Huang Y. et al. provided are bias, misinformation, intellectual property, and transparency, which are systematically mapped after successfully reviewing 39 different studies related to LLMs. Their conclusions point to the fact that technical solutions should be reinforced by policy, education, and human control (Huang Y. et al., 2025). The unsupervised sentiment and topic analysis of close to 24,000 Reddit posts conducted by Xu Z. et al. showed that the perception of ChatGPT remained largely positive based on results obtained after sentiment analysis depicting nearly 61.6% of the posts and comments were in favour of ChatGPT, but there is still a concern regarding the usability and ethical threats (Xu Z. et al., 2025). Garrido A. P. and Barra E. emphasized that education required explainable and fair transformer models on matters where vulnerable populations are involved; this is after they had critically analysed 41 articles (Garrido A. P. and Barra E., 2025). Luo X. et al. reported a change in attitude of the users of Reddit towards skepticism and then cautious optimism on Reddit discussions available for January 2024 using the Reddit API. This study focused on the query “ChatGPT” and “Therapy” and concluded with the conditions of overcoming the ethical and usability issues successfully (Luo X. et al., 2025). Even with their insights, most of the ethical studies fail to take into consideration the variability of contexts and fail to evaluate how ethical norms change over time in online discussions.

In methodology, topic modelling, lexicon-based sentiment, or simple classifiers are used in most ethical studies that are conducted on moderate-sized datasets. They do not often use a single pipeline to combine fine-grained labels of ethical stance with sentiment and creativity. A number of those reviews highlight risks like bias and opaqueness but fail to operationalise such issues on a large-scale quantitative basis. This gap is filled in my work, which considers ethical stance as a quantifiable attribute of every post, where zero-shot BART classification is done on the entire Reddit corpus.

3.4 Social Network Analysis and Reddit

The social network analysis has made important contributions to the dynamics of GenAI discourses. The circulation of the creative and ethical norms in the structures of the Reddit is tracked in a subtle form as demonstrated by Mentzer K. et al. in their study (Mentzer K. et al., 2024). Nedungadi P. et al. combined graph neural networks with transformer text mining, which showed that there are dynamic patterns of influence, community formation, and topic spread (Nedungadi P. et al., 2022). Pandey G. et al. presented meta-analysis data in support of scalable models, which follow creativity outputs, norm adoption, and controversies based on the data of Reddit (Pandey G. et al., 2025). Tunca S. used the same approaches which included sentiment analysis using VADER orchestrated with BERTopic which acted as a transformer-based topic model to neurodivergent communities and revealed previously concealed opinions and emotional patterns on a dataset comprising of 238,100 Reddit posts and comments (Tunca S., 2025). Weaknesses in this literature are that there are difficulties in multimodal and multi-cultural data and that network building algorithms can be biased.

Available studies of Reddit networks normally recreate entire reply trees or operate on subreddit-scale graphs with full metadata. Influence, community structure, and topic diffusion over time are used to analyse graphs by algorithm (or neural networks) applied to them. Conversely, sparse comment identifiers and blank subreddit fields of the majority of my records mean that a lightweight

co-occurrence network is the only possible alternative. This limitation provokes the temporal and user-level analysis, but it still depicts fundamental educational tags and subjects in the discussion.

3.5 Obstacles of Prejudice, Equity and Social Effects

The industry is marked by continued issues of algorithmic bias, lack of transparency of the models used, adversarial susceptibility and incapacity to address various language and cultural settings, as Masciari E. et al. describe in their article about AI-based recommendation systems and ethical aspects (Masciari E. et al., 2024). In the opinion of Ray P. P. in his detailed review, ChatGPT as a chatbot technology continues to have historically old issues: bias in algorithms, the inability to explain, adversarial vulnerability, and inefficiency in processing multimodal and multi-linguistic and culturally diverse content (Ray P. P., 2023). Kowalkiewicz M. in his book *The Economy of Algorithms: AI and the Rise of the Digital Minions* highlighted the necessity to address the current gap between machine-generated output and human ethical evaluation with the help of human supervision and cooperation with other fields (Kowalkiewicz M., 2025). Lv Z. in his article emphasized the necessity to have common standards of benchmarking, transparency, and explainable models that are energy-efficient in order to promote responsible ideation in creative process under GenAI (Lv Z., 2023). Both the articles by Chiarello F. et al. and Nedungadi P. et al. have suggested strong model development and enhanced mixed-method methods that integrate qualitative information with extensive quantitative studies as the way of the future. The new agreement indicates that the most responsible way to implement GenAI is through synergy between strong algorithms and active participation of human beings, where AI capabilities in recognizing patterns are used, and human creativity and ethical principles are grounded.

3.6 Future Projections and Unresolved Issues

With increasing functionality of transformers and social network analysis tools, literature has started to discuss open issues of transparency, fairness, and domain adaptability, Chiarello et al. encourage researchers to pay attention to model robustness and adversarial defense (Chiarello F. et al., 2024), whereas Nedungadi P. and others demand more detailed mixed-method techniques that combine qualitative knowledge with high-scale quantitative references (Nedungadi P. et al., 2022). This renewed literature review provides a more refined and critical insight into the existing literature, which provides a strong framework of theoretical and methodological basis of this study. It determines existing advantages and major weaknesses, particularly with regard to the evaluation of user creativity and ethical decision-making with regard to the actual interactions of GenAI.

In general, the literature that is reviewed supports the claim that transformer models enhance semantic representations, although much of the research includes comparisons of accuracy on benchmark tasks only. They do not often measure the level of creativity or moral consideration that AI use deems in the daily practice, particularly in academic communities. Innovative studies are biased to short-term interactions of experiments and ignore long term, everyday applications of GenAI in education. Bias and misuse are listed as the dangers, though ethical studies rarely relate the ethical categories to the language patterns that are present in large social media corpora. My work addresses these gaps by integrating sentiment, ethical position, creativity measurements, and network structure into a single, unified, pipeline on a dataset of huge size with a focus on education and with the use of Reddit.

4 Research Methodology

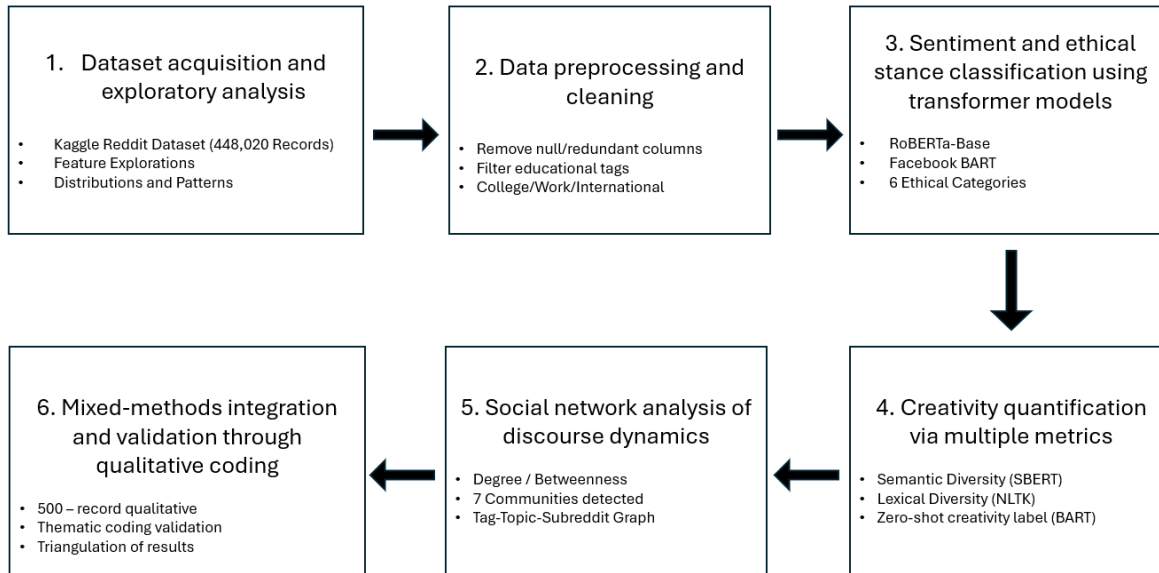


Figure 1 - Study Design

The study design as shown in Figure 1, applied in the research is a mixed-methods one that combines quantitative transformer-based natural language processing (NLP) analysis using RoBERTa, Sentence-BERT (MiniLM) and BART with qualitative thematic coding and social network analysis to investigate the impact of the use of generative AI tools on user creativity and ethical decision-making in educational Reddit communities. The pipeline of the methodology includes six consecutive steps: (1) Dataset acquisition and exploratory analysis, (2) Data preprocessing and cleaning, (3) Sentiment and ethical stance classification with the help of transformer models, (4) Quantifying creativity with the help of multiple metrics, (5) Social network analysis of discourse dynamics, and (6) Integration and validation of the results with the help of qualitative coding. This research method provides a coherent fit to the research question and makes it possible to thoroughly examine semantic, ethical, and network aspects of the Reddit discussion about ChatGPT.

4.1 Data Collection and Exploratory Analysis:

The data used in the research is the “How Students Use ChatGPT Filtered Reddit Dataset”, which was taken from Kaggle ([Link](#)) ([License](#)). The initial dataset consists of 448,020 posts and comments about education and ChatGPT published on Reddit collected over the period from December 2022 to mid-2025, which are filtered in terms of educational and ChatGPT-related materials, and it is a discussion in the various communities on the subreddit. The dataset was filtered to include natural user speech on ChatGPT in educational, social, and professional settings, which is appropriate to analyse the true student attitudes towards the use of generative AI. The dataset was formed using a combination of various public Reddit datasets on Kaggle and filtering them with the help of specific keywords as mentioned by the dataset publisher Mr. Muhammad Faizan.

The initial dataset (448,020 records) had ten variables and was further analysed by loading the dataset and naming the data frame as ‘df’ using pandas library. These variables include ‘subreddit’ (Reddit community identifier), ‘body’ (text content of a post/comment), ‘controversiality’ (binary flag which shows controversial posts), ‘score’ (number of upvotes-downvotes), ‘topic’ (pre-assigned topic identifier), and ‘tag’ (categorical label College, Work, International, etc.). Also, there were redundant fields, such as ‘Unnamed: 0’ (row index), ‘comment id’ and ‘comment parent id’ (identifiers of comment hierarchies to use in network analysis), and comment body (duplicate text field). It is important to note that significant amounts of the missing values existed: 435,818 records (97.3% of the total) did not have subreddit information, and 12,202 records (2.7% of the total) did not have topic and tag assignments. This study discarded the variables whose values were either structurally missing or not informative in any way (Unnamed: 0, comment id, comment parent id and comment body) because they were practically all nulls and did not have any impact on the main NLP tasks. On the analytical variables of topic and tag the study considered complete in the retained subset df updated and viewed the large missing field of subreddit as missing, instead of trying to impute it, so analyses

and visualisations which used subreddit information were explicitly identified as relying on the 2.7% of records containing such metadata.

According to EDA we performed on our initial data frame 'df', the following distributional properties of data frame 'df' were discovered which enlightened us to create clean data frame 'df_updated':

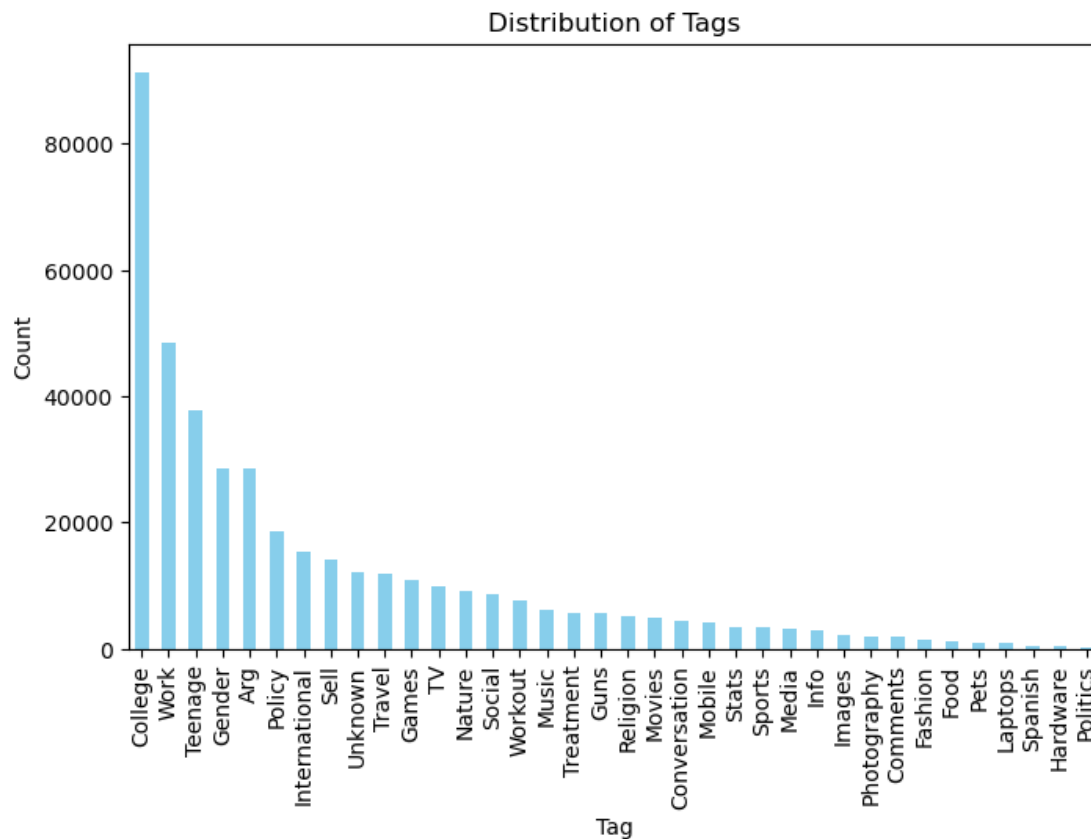


Figure 2 - Tag Distribution for ChatGPT-related Reddit posts

- Tag Distribution for ChatGPT-related Reddit posts (Figure 2): The most frequent tag was college, which was used as 57.1% of all tag records ($n = 101,291$). Work was represented by 13.8% ($n = 24,597$), International by 9.2% ($n = 16,044$), and Social by 7.6% ($n = 13,262$) with minor proportions of Stats (4.1%), Treatment (3.6%), and Conversation (2.2%). This clustering of the tags around education stimulated the choice to keep only the posts in which tags were obviously education or study related.

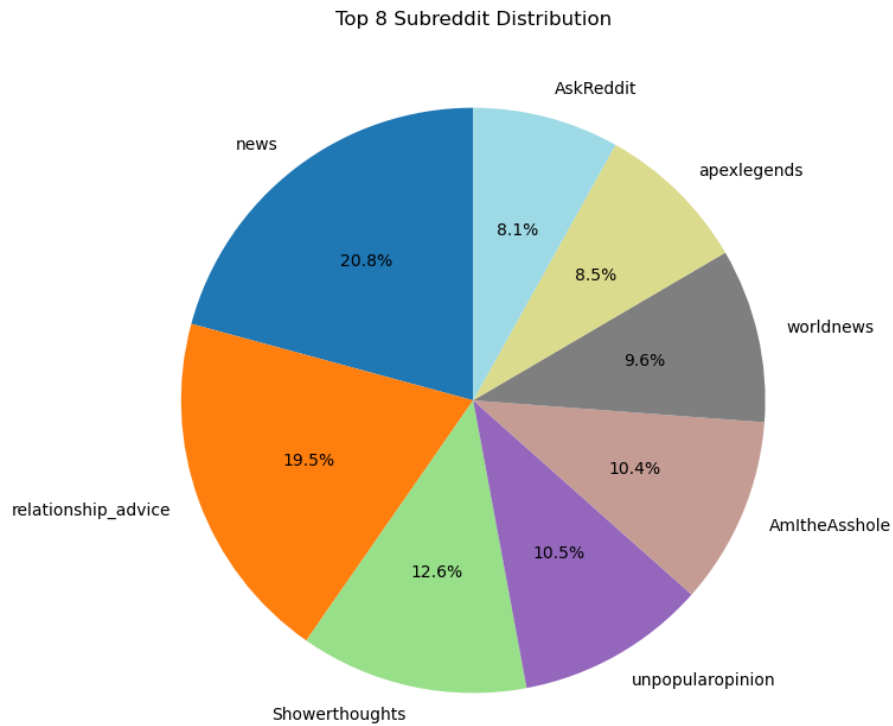


Figure 3 - Distribution of Top-subreddits

- Distribution of Top-subreddits with considered dataset (Figure 3): Of the records that contained subreddit metadata (2.7% of the dataset), news was the most common subreddit (20.8%), with relationship advice (19.5%) and AskReddit (8.1%), and a long tail of other communities. The high number of lost values of subreddits shows that the ChatGPT discussion did not take place in a limited number of forums but was spread among numerous communities.

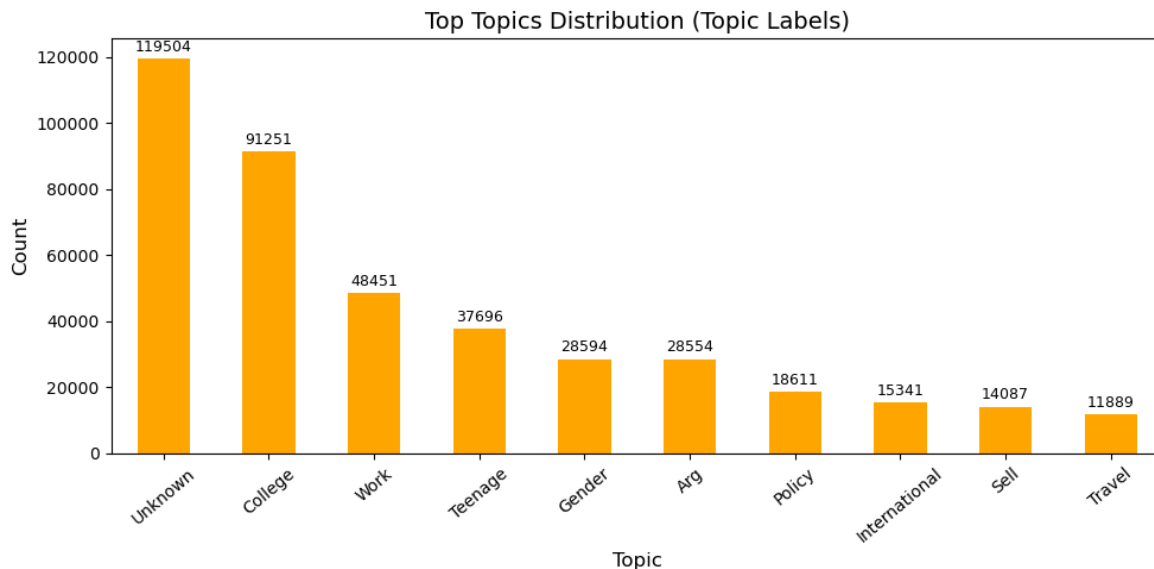


Figure 4 - Topic Distribution Bar Chart

- Topic Distribution Bar Chart (Figure 4): By the number of records that contain the topic named with the relevant tag name, the largest group is Unknown topics (n = 119,504), then there is the College (n = 91,251), Work (n = 48,451), Teenage (n = 37,696), Gender (n = 28,594), Arg (n = 28,554), Policy (n = 18,611), International (n = 15,341), Sell (n = 14,087) and Travel (n = 11,889). This distribution indicates that, following the massive residual 'Unknown' group, the most represented labelled topics are education- and youth-related topics like College, Work, Teenage and International, which plays into the high educational orientation of the data.

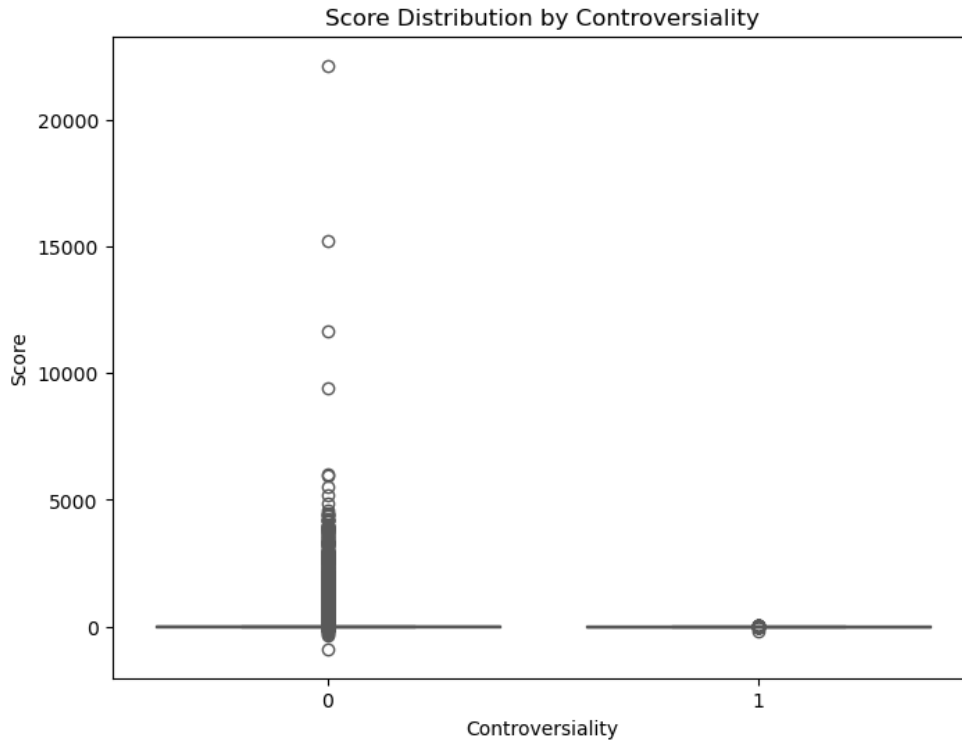


Figure 5 - Score vs Controversiality Boxplot

- **Score vs Controversiality Boxplot (Figure 5):** The score-controversiality boxplot indicates that the bulk of posts has low to moderate scores with an interquartile range of 1 to 3 and a median of 2, and few outliers have very high scores (up to 22,159). Controversial and non-controversial posts may receive high scores, which implies that controversy is not the sole motivator of participation in such debates.

4.2 Data Preprocessing and Cleaning

The preprocessing pipeline consisted of three steps, namely:

Step 1: Selection and Redundancy of Features. Columns with a large number of null values, or with the null value, were dropped, namely: Unnamed: 0 (100% null), comment-id, comment-parent-id and comment-body (98.9% null values). The body column (post/comment text) was kept as the main element of NLP analysis, whereas the score and controversiality measures were kept to the contextual analysis.

Step 2: Relevance Tag-Based Filtering. The screening had been conducted to retain only those posts that were tagged with educationally relevant context: College, Work, International, Social, Stats, Treatment, Conversation, Study, Education and School. This filtering removed 270,535 records (60.4 percent of the original 448,020 records) and included 177,485 records (39.6 percent), on which all further sentiment and ethical stance and creativity and network analysis was performed. The selection of educationally oriented tags helped to retain the final dataset of student and professional interactions with ChatGPT in a relevant context.

Step 3: Final Data Construction. After cleaning, the cleaned data contained 177,485 records (39.6% retention rate, 60.4% data loss) and was stored as a new data frame called as 'df_updated'. The loss incorporates a deliberate quality and relevance filtering; the deleted records were more off topic than corrupted. The final processed data included six variables (subreddit (community identifier), body (text content), controversiality, score, topic, and tag (educational category)).

4.3 Transformer Model Based Quantitative Analysis

Sentiment Classification analysis is done using two transformer-based models using the Hugging Face transformers library, implemented using a batched pipeline model. RoBERTa-Base (Twitter Sentiment) was considered as the first model. The cardiffnlp/twitter-roberta-sentiment-base-most-recent model, a 3-class sentiment RoBERTa fine-tuned on 124M tweets, was the main sentiment classifier. RoBERTa is an extension of the architecture of BERT with robustness optimization such as dynamic masking and larger batch training. The choice of this model was based on its effectiveness in

working with social media texts and being sensitive to informal linguistic features of the conversation on Reddit. Secondly, Facebook BART (Ethical Stance) is considered as the second model. Although it was mainly applied to zero-shot ethical classification, BART encoder component also provided contextual knowledge about subtle language.

The 177,485 number of posts/comments was segmented into 16 records batches to optimise the computational efficiency. To run the Hugging Face pipeline on the batch, the raw text in the body field was run through a Hugging Face pipeline('sentiment-analysis', model=cardiffnlp/twitter-roberta-base-sentiment-latest) which returned, per record, a prediction of a sentiment label, either positive, negative or neutral, along with the confidence score of the model between 0 and 1. These two outputs were then rewritten into the working data frame in the form of new columns, sentiment and sentiment_score, that is, the dataset did not initially have sentiment annotations, but this transformer-based inference step created them completely.

Ethical Stance Classification aims to identify ethical perspectives on a specific issue that may also be linked to the concept of cultural competence. This method is going to determine the ethical view on a given problem that can also be associated with the notion of cultural competence. The determination of ethical stance was through zero-shot classification, which allows making a classification without training data that are specific to the task by utilizing the knowledge of large language models. The approach involved Facebook BART-Large-MNLI architecture. Zero-shot ethical stance classification using the facebook/bart-large-mnli model, which was trained on the Multi-Genre Natural Language Inference (MNLI) corpus, was used. BART is an encoder-based architecture of BERT with a transformer decoder capable of doing flexible input-output mapping. The model was trained on 433,000+ premise-hypothesis pairs of 433,000 premises, 433,000 hypotheses, to infer arbitrary texts to predefined ethical categories using natural language. Labels and Classification of Candidates include Seven categories of ethical stance were identified and defined by previous literature: prosocial, neutral, harmful, ethical, unethical, moral dilemma and unethical. Similarity scores between the input text and the candidate label, which were represented as a text stating a stance of this or that label, were used to classify texts. Each record was rated with the highest score label.

4.4 Creativity Quantification

Three complementary measures that combine quantitative and qualitative zero-shot classification measures are used to measure creativity:

Metric 1: Semantic Diversity (Sentence-BERT). Semantic diversity measures the novelty of conceptual material. The entire 177485 posts were encoded to 384-dimensional embeddings using the all-MiniLM-L6-v2 Sentence-BERT model. A sample group of 1,000 random posts was used to come up with a baseline. Cosine similarity was then calculated between each record and the reference set, and diversity scores were obtained as $1 - \text{mean}(\text{similarity})$ with a range of 0.75-1.05. The more the value, the more the conceptual novelty.

Metric 2: Lexical Diversity (Token-Based). Lexical diversity determines the vocabulary richness, in terms of tokens (unique tokens) to the summation of unique tokens:

$\text{Diversity} = (\text{unique_tokens}) / (\text{total_tokens})$.

Records that were 0.0-1.0 lexical diversity used NLTK tokenization.

Metric 3: Zero-Shot Creativity Label (Qualitative Indicator - BART). Similar to the classification of ethical stance, there were five categories of predefined creativity, i.e.: highly creative, original, routine, informative, humorous. The BART-MNLI model and zero-shot approach were used to label each record with a creativity label.

4.5 Social Network Analysis

The hierarchical nature of Reddit, in which the comments respond to parent comments in threads, can be used to network analyse discourse dynamics. Directed post-comment graphs were built using the Python package 'NetworkX', which is used to build, manipulate and analyse complex networks. NetworkX offers graph centrality measures (metrics) algorithms, community detection algorithms, and visualization algorithms of network topology.

The process of construction of the network included:

Edge List Formation: Comment parent to child relationships were obtained by the comment id and comment parent id fields, with an ordered relationship between a parent and a child a parent-to-child comment relationship.

There were also co-occurrence relationships: the edges linked tags to topics, tags to subreddits, and topics to subreddits, which created a multipartite network of 15 nodes that represented 8 categories and 7 subreddits.

In terms of Graph Metrics Computation, there were two centrality measures that measured the importance of nodes:

Degree Centrality: The ratio of the number of nodes that are directly linked to a particular node indicates the topic or community that is influential.

Betweenness Centrality: It is a metric of the number of shortest paths that go through a node, which is used to determine information or norm hubs. The nodes that have high betweenness centrality serve as intermediaries between discourse clusters that would otherwise be far apart.

Community Detection highlights the greedy modularity optimisation identified seven communities in the tag topic subreddit network with the community sizes ranging between 23 nodes. They were represented by thematically related clusters of discourse (e.g., educational settings, professional practices, ethics).

4.6 Mixed Methods Augmentation and Integration

To overcome possible automated classifier biases and gain better interpretability, quantitative results were tested using qualitative thematic analysis. A stratified random sample of 500 records was chosen and this showed proportional samples in regard to sentiment (neutral: 50.6%, $n = 253$; positive: 31.8%, $n = 159$; negative: 17.6%, $n = 88$), ethical stance (prosocial, neutral, harmful, moral dilemma, unethical, ethical) and level of creativity (informative, original, humorous, routine and highly creative). This qualitative sample was thematically coded assuming: (1) close reading of the text of the post, (2) comparison with automatically-assigned labels to evaluate agreement, (3) finding discrepancies in coding and the reasons behind them (e.g., sarcasm, subtle language, cultural reference), and (4) recording contextual factors that can influence sentiment and ethics that might be missed by automated models.

4.7 Justification and Limitations

The mixed-methods fusion has the scalability of quantitative (177,485 records in a few minutes), and the depth of qualitative (stratified thematic validation of classifier results). The visualisation of social network analysis exposes the phenomena at the community level that cannot be seen in the text-only analysis and shows the propagation of norms and collaborative discourse structures.

- **Model Bias:** The transformer models have an inheritance of biases of training data. It is possible that RoBERTa has not been trained on the linguistic norms of Reddit. The MNLI training provided by BART might not be a complete reflection of ethical dilemma in AI practices. Classifiers can be systematically misclassifying sarcasm, negation or domain specific terms.
- **Obvious Ethical words:** Zero-shot classifiers can have a difficult time identifying implicit ethical reasoning, hedging the words (I think, maybe), or ethical subtlety in context. A particular record in the dataset could raise concerning issues with reliance on AI which may be labelled as neutral, even though it is ethically worried.
- **Lacking Data Completeness:** Non-textual data (images, videos) are not covered in the dataset, and the metadata (the time when the posts were published) is not presented, which restricts the accuracy of longitudinal analysis. The information on subreddits is not available in 97.3 percent of records, making it impossible to compare it on a subreddit level.
- **Sampling Bias:** The Kaggle dataset was narrowed down to include only students and ChatGPT, which may not represent a general conversation on the topic of generative AI or the non-student experience.
- **Network Sparsity:** The post-comment network is not fully reconstructed because comment-id/parent-id loss in the preprocessing part of the network analysis restricts the granularity of the network analysis.

5 Implementation

The last implementation phase aimed at converting the cleaned Reddit data into an analytically useful final form via a series of NLP and network-analysis pipelines, which were written in Python. The task was executed in Jupyter using Python 3, with data manipulations conducted with the help of pandas and NumPy, plotting conducted with the help of matplotlib and seaborn, and model and analysis conducted with the help of scikit-learn, NLTK, sentence-transformers, transformers, and NetworkX.

To start with, the filtered data frame, `df_updated` (177,485 records; six columns: subreddit, body, controversiality, score, topic, tag), was transferred to `df_proc` as the working table of all downstream processes. The `cardiffnlp/twitter-roberta-base-sentiment-latest` model of Hugging Face sentiment-analysis was then run in batches of 16 texts, which produced a discrete sentiment label (positive, neutral, negative) and a confidence score that was stored in the new columns `sentiment` and `sentiment_score`.

Then, the classifier of the ethical stance (zero-shot) was applied to a `facebook/bart-large-mnli` model, and six candidate labels were used: prosocial, neutral, harmful, ethical, unethical and moral dilemma. Ethical stance label of each record was generated with batched inference (batch size 8), and the in-between results were stored in `reddit` with `sentiment ethics.csv` to prevent the rerun of the lengthy classification jobs.

The application of creativity quantification was done in three sections. A semantic-diversity score (`creativity_semdiv`) of each text was computed (sentence- Bert- Embeddings (all-MiniLM-L6-v2) as one minus the average cosine similarity to a 1,000-post reference set. Looking at the NLTK tokenisation generated a score of lexical diversity (`lexical_diversity`) which is the ratio of unique tokens to total tokens. Lastly, a second BART zero-shot pipeline was used to label each post with a label out of high creative, original, routine, informative, humorous, which was stored as `creativity label`.

To analyse the network, the NetworkX was constructed on simple undirected graph `G` of co-occurring tag, topic, and subreddit values in `df updated`. Degree centrality and community detection were calculated, and a smaller subgraph of the most centralized nodes was plotted as the Top Tag/Topic/Subreddit Network that was used in the Results and Discussion sections.

6 Evaluation of Results

The results of the analytical pipeline are assessed in this section on four case studies that are interconnected. Case Study 1 explores the topic of emotional framing of ChatGPT by students through sentiment distributions and tag level trends. Case Study 2 deals with ethical position, where the discourse is either prosocial or neutral, or even harmful. Case Study 3 measures creativity with a combination of semantic and lexical diversity and zero-shot creativity nomenclature. The Case Study 4 explores the discourse organization based on the use of a co-occurrence network of tags, topics, and subreddits, where discourse focuses and which themes serve as centres. Collectively, these case studies show that sentiment, ethics, creativity, and network structure can be used together to describe the student discussion about generative AI on Reddit.

6.1 Case Study 1 - Sentiment Analysis

Case Study 1 concentrates on the role of the emotional positioning of ChatGPT among students in the learning settings. It is expected to measure the total proportion of positive, negative, and neutral sentiment and to observe how these emotional tones change by educational tags of College, Work, and International. This is used to decide on whether ChatGPT is mainly felt as a useful helper, a frustrating tool, or a neutral tool to do routine academic work.

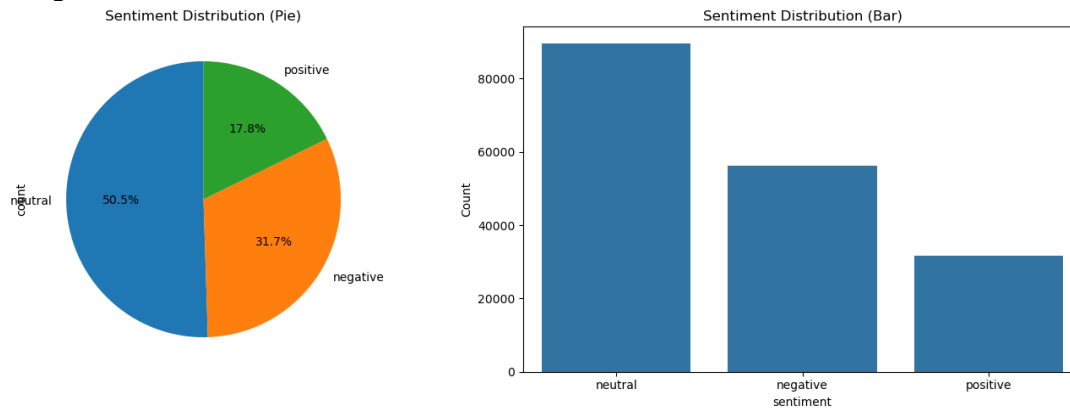


Figure 6 - Sentiment-Distribution plot

The sentiment-distribution plot in Figure 6 indicates that there are more neutral posts than any other type of post as 50.5% of all records are neutral and 31.7% and 17.8% are negative and positive respectively. This suggests that the discourse of students using ChatGPT within the educational setting is usually descriptive or informative as opposed to highly polarised. Methodologically, the three-way divide is in line with the RoBERTa Twitter model design, which is optimised to informal social-media language and thus well-adapted to Reddit text. The bar plot affirms the presence of large numbers of absolute counts in the three classes, which is an indicator of strong comparisons among sentiment categories.

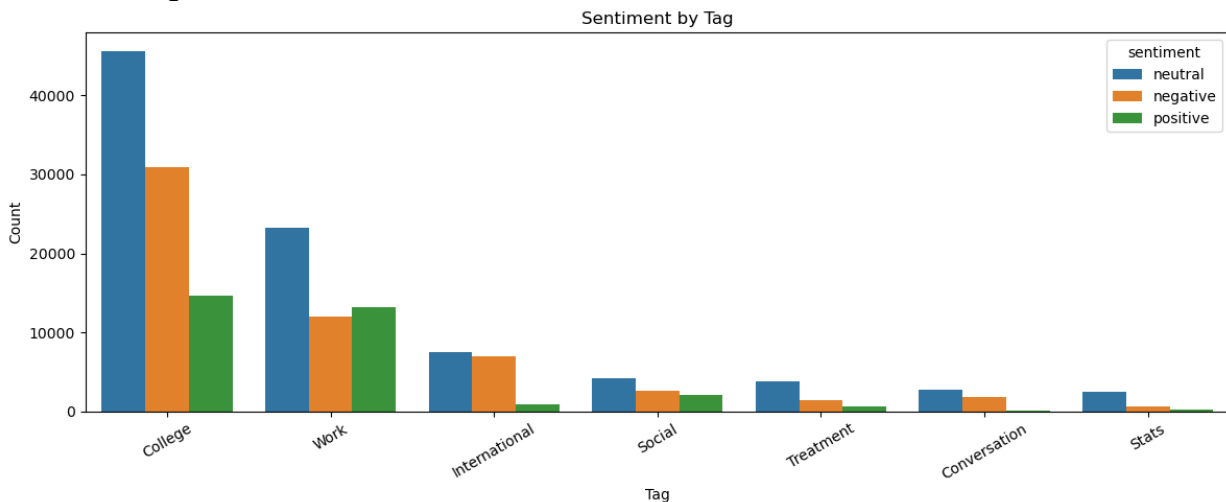


Figure 7 - Sentiment by Tag Bar chart

According to the Sentiment by Bar Tag chart in Figure 7, there is a significant difference in the sentiment among the educational tags like College, Work, International, Social, Treatment, Conversation and Stats. College and Work posts demonstrate large quantities of neutral and positive attitude and a relatively small number of negative postings indicating that numerous students report to use ChatGPT as pragmatic assistance (e.g., clarification, drafts) and not as a source of frustration. Other tags, such as Conversation and Social, have less positive or rather negative proportions and they could indicate the discussions and disagreements or the critical commentary on the use of AI. These tag-level distributions demonstrate that ChatGPT is perceived differently in different educational contexts, which directly respond to the research goal of mapping the sentiment in various cases of student use.

6.2 Case Study 2 - Ethical Stance Analysis

Case Study 2 examines an ethical aspect of Reddit posts on ChatGPT. It is aimed at quantifying the frequency of posts with prosocial or responsible attitudes, ethically neutral talk, or potentially harmful behaviours including cheating, over-reliance or misuse. This case study evaluates the relationship between emotionally charged posts and those that have ethical risks by cross tabulating the ethical stance and the sentiment.

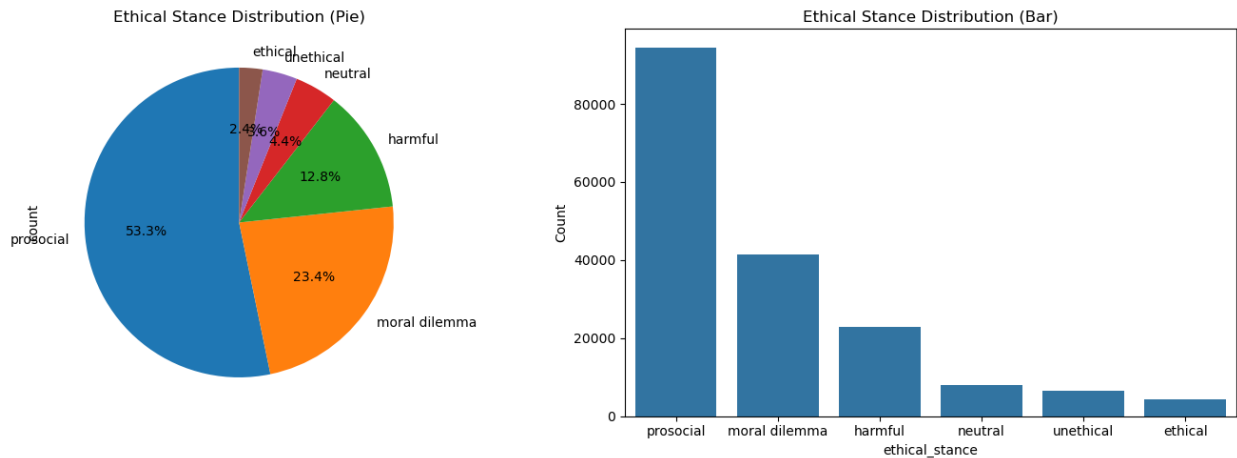


Figure 8 - Ethical-Stance Distribution

According to the ethical-stance distribution in Figure 8, the majority of the classes are prosocial (53.3%), then moral dilemma (23.4%), and harmful (12.8%), and the rest of the share is occupied by the moral neutral, ethical, and unethical labels. This trend implies that the majority of the postings are either constructive, supportive or ethically neutral, with a minor but significant part of the posts mentioning potentially harmful or ethically problematic issues. Since BART MNL model can infer labels through the inference of natural language, instead of the target training, these proportions give a high-level index of the frequency of ethical reasoning or risk manifestation in the discourse.

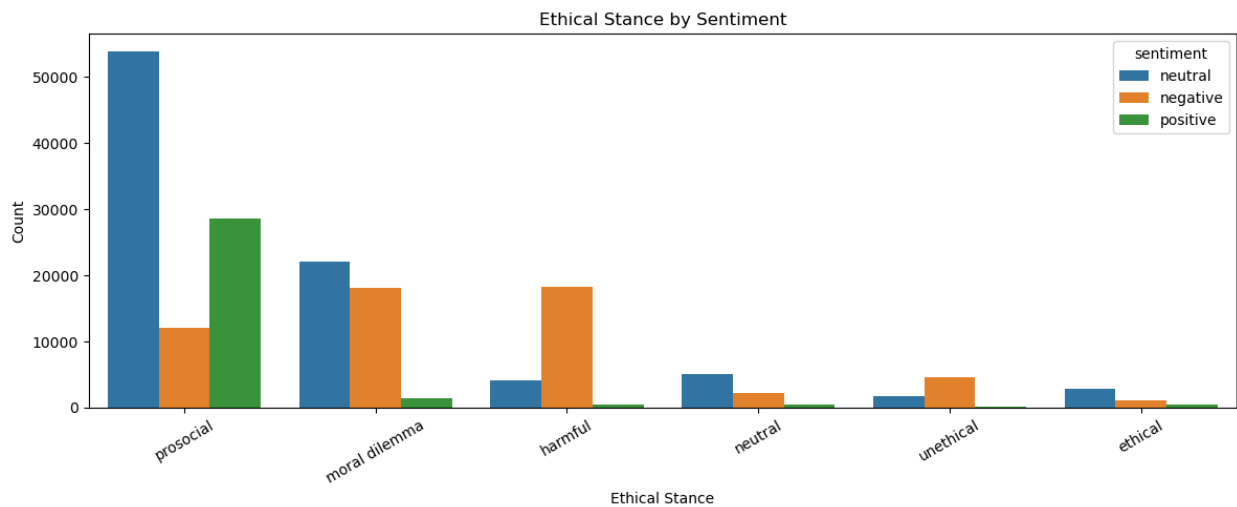


Figure 9 - Ethical Stance by Sentiment Bar Chart

The chart of Ethical Stance by Sentiment in Figure 9 provides a more in-depth diagnostic. Posts that are prosocial are likely to be co-occurring with positive sentiment, whereas harmful and moral-dilemma posts are associated with negative sentiment as predicted when students give descriptions of cheating, dependency, or misuse. The neutral sentiment is propagated in multiple categories of ethics, which highlights that it is not always that ethically relevant content is emotion-charged. This cross-tabulation assists in affirming the consistency of sentiment and ethics pipelines which indicate that they represent the opposite sides of the same discussions.

6.3 Case Study 3 - Creativity Quantification

Case Study 3 focuses on the extent to which creative or routine student talk about ChatGPT is. The case study will differentiate between the posts where information is merely requested and those where original, humorous, or highly creative interaction with AI takes place using semantic diversity, lexical diversity, and transformer labels of creativity. This explains why ChatGPT is being utilized as a creativity engine or more as an informational assistant.

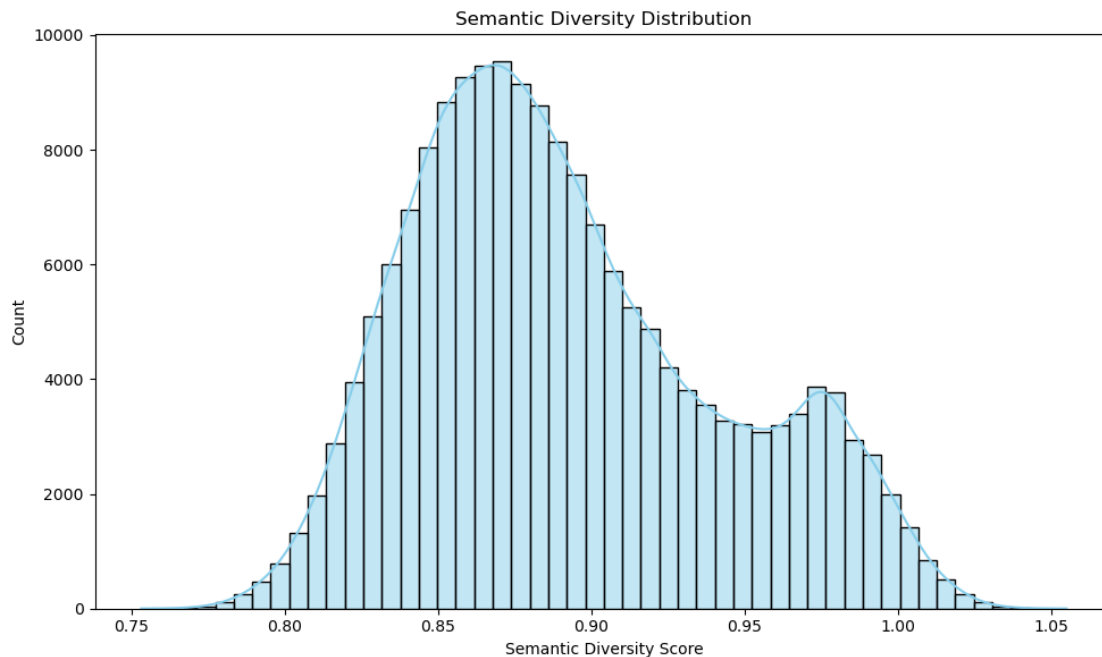


Figure 10 - Semantic Diversity Histogram

The semantic-diversity histogram in Figure 10 is unimodal whereby about 82% of the posts (n 145,500 out of 177,485) have semantic diversity between 0.85 and 0.95, which means that most of the texts are neither highly similar to the 1,000-post reference set, yet still share a similar semantic range. It implies that students normally present different but connected situations and queries-which is in line with classroom or study-support application-and a smaller tail of posts with diversity scores bigger than 0.95 will be truly unusual or an outlier discussion. This implies that students often introduce variations of related issues or situations as opposed to completely new ones, which is reasonable in a classroom question and answer environment. Meanwhile, the tail of the higher diversity values indicates honestly bizarre conversations, which is consistent with the interest in the outlier creative applications of ChatGPT.

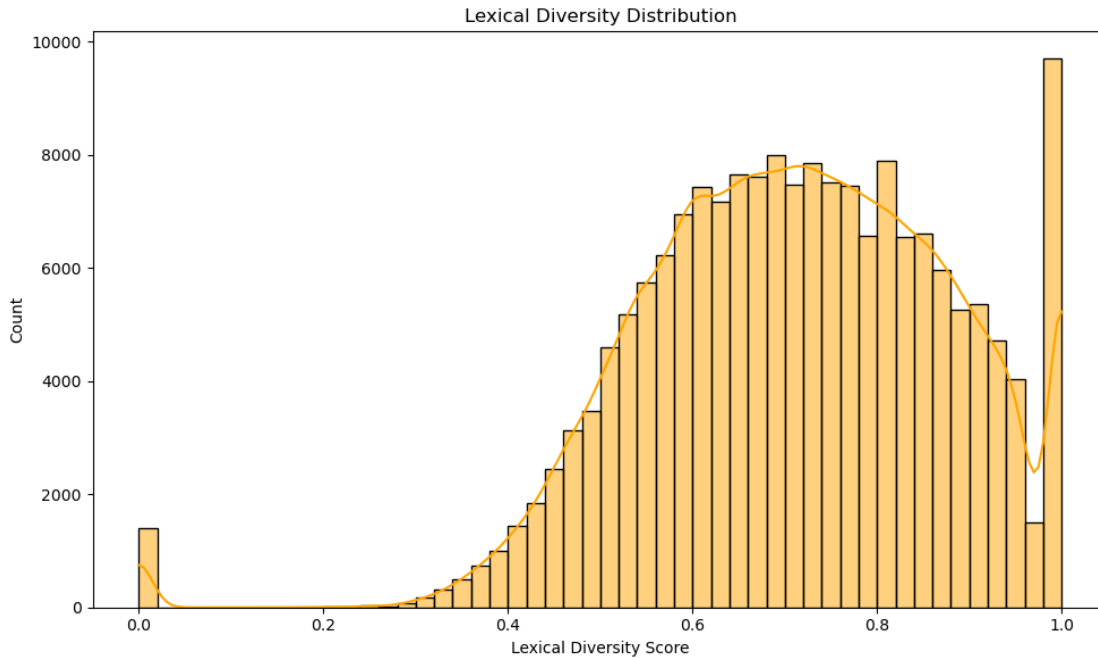


Figure 11 - Lexical Diversity Histogram

The distribution of lexical diversity as seen in Figure 11 is unimodal as well with approximately 69% of posts (n 122,500) having lexical-diversity values between 0.4 and 0.7, and 31% posts with a value above 1.5, indicating a middle range where users balance between repeating core academic vocabulary (e.g., assignment, essay, exam) and some variation. The percentage of posts with a lower score than 0.2 or a higher score than 0.8 is minimal, which implies that the frequent use of very formulaic or very diverse words is not common in this sample of educational Reddit. Since lexical diversity is calculated on a per-post basis, low values in very short comments will be anticipated and do not necessarily mean low creativity, therefore, a combination of this measure with semantic diversity and the zero-shot creativity labels will offer more insightful data.

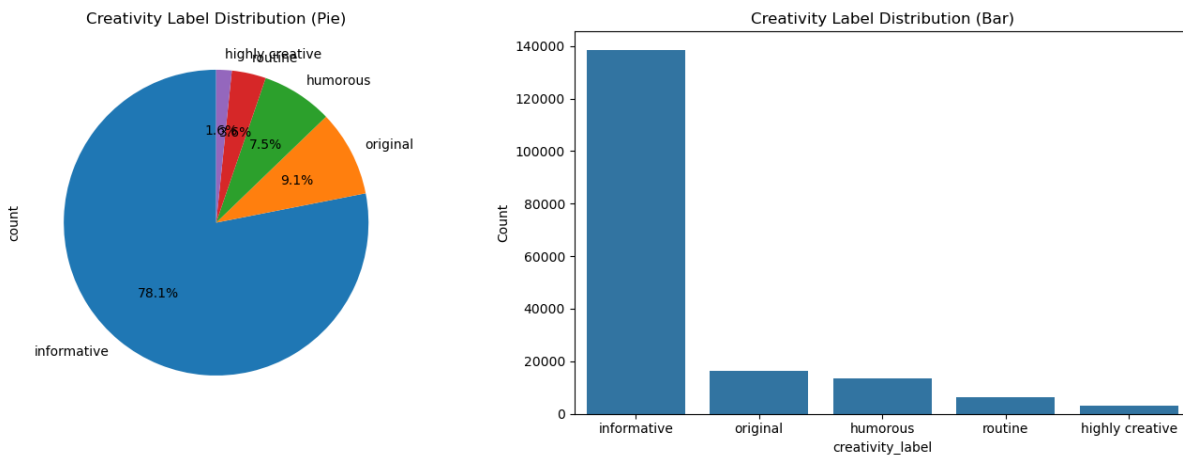


Figure 12 - Creativity Label Plots

The creativity-label plots in Figure 12 indicate that 78.1% of posts are ‘informative’ while ‘original’, ‘humorous’, ‘routine’, and ‘highly creative labels’ cover 9.1%, 7.5%, 3.6%, and 1.6% respectively. This is a strong indication of the meaning that students primarily use ChatGPT in an informational or problem-solving context, specifically, they request explanations, clarifications, or worked examples, but not an open-ended and creative approach. The non-trivial ratio of funny and original posts suggests that there are users who play around or improvise but this is not the primary focus. Academically, these allocations support the argument that GenAI is not yet a tool that replaces creativity among students, but rather as one that assists them.

6.4 Case Study 4 - Network Analysis

Case Study 4 discusses the relational form of discourse through the modelling of co-occurrences of tags, topics and subreddits as a network. It is proposed to find the key educational scenarios, like College or Work, and learn to relate them to the topics and communities in which norms of ChatGPT use are negotiated. This network image supplements the distributional analyses that have been made before by demonstrating where people are talking and where the nodes serve as centres of influence.

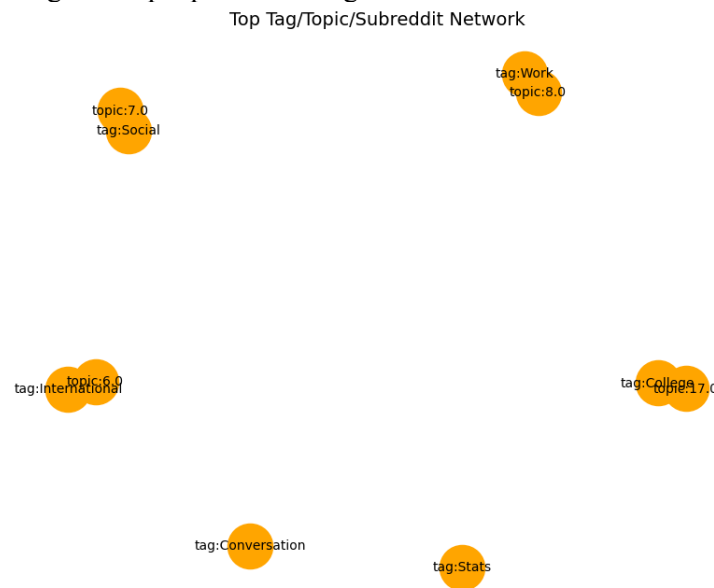


Figure 13 - NetworkX Graph Plot

The NetworkX graph plot in Figure 13 indicates that there is a small yet informative co-occurrence graph of 15 nodes and 8 edges between tags, topics and subreddits. Degree centrality is used to identify tag:Conversation with the highest connection (0.143), next are tag:College, tag:Work, tag:International, tag:Social, tag:Stats, and 17.0, 8.0, 7.0 and 6.0 topics (all with 0.071). It suggests that discussion strands that connect several tags and subjects serve as forums where the standards regarding the use of ChatGPT are being bargained. Greedy modularity community detection identifies seven small communities of size 2-3 that represent closely knit clusters that include College-topic 17.0 or Work-topic 8.0. The visualisation of the “Top Tag/Topic/Subreddit Network” brings out the linkages between educational tags and particular discussion themes and subreddits, providing a relational dimension to the previous distributional results. Combined, these metrics of network data demonstrate that the discussion of ChatGPT is not evenly distributed but clustered around several major educational scenarios, which can be used to intervene in high-centrality communities (e.g., through guidelines or training).

In all four case studies, the usage of ChatGPT by students is largely informative and prosocial, and with some areas of concern related to harmful or ethically questionable usage and limited yet present creative exploration.

7 Discussion

This paper aimed to examine how creativity and ethical choices in educational Reddit communities are manifested in the use of ChatGPT by students based on transformer-based NLP and social network analysis of 177,485 posts collected over the period from December 2022 to mid-2025. The results of the sentiment analysis, ethical stance, creativity metrics, and network structure combined together point to the conclusion that ChatGPT is embedded as an informational and support tool mainly, and more creative and ethically problematic applications are more limited but not absent.

The prevalence of neutral and positive feeling (50.5 and 31.7 percent, respectively) indicates that the majority of student experiences with ChatGPT are defined by descriptive, pragmatic, or appreciative intonation instead of polarised excitement and aggression. This is in line with the recent literature that has placed GenAI as an extension of cognitive, and not a radical disruptor in the daily learning activities. Meanwhile, the nearly one-fifth negative sentiment indicates the existence of certain areas of frustration, concern, or disappointment that could be of greater interest, especially in the tags where the fairness, cheating, or workload discussion seems to be more prevalent.

The prevalence of neutral and prosocial posts in my outcomes confirms previous research that people have more positive or pragmatic views about ChatGPT than alarmist. Most of the positive sentiment on Reddit is also reported by Xu et al., although there was no particular concern with the ethical position as in this study. My results of creativity, in which posts tend to be informative and those of a high level of creativity are uncommon, conform to Desdevises and Jackson who characterize GenAI as fluent and conventional. Nevertheless, when I apply semantic and lexical diversity on a large scale, I demonstrate that students remain exposed to a broad variety of related situations instead of merely repeating templates. This network analysis resonates with Mentzer and others in that GenAI debates are spread through numerous communities, though my emphasis on education tags relates to how contexts related to college and work serve as centres of play.

An ethical stance analysis introduces a second dimension because it shows that more than half of the posts are prosocial and approximately a quarter of them are neutral, whereas a smaller yet significant portion are harmful, unethical, or presented as moral dilemmas. This distribution indicates that students do not simply use ChatGPT as an indifferent utility but will often interact with appropriately used, fairly and ethically constructed and scholarly-level questions. The presence of harmful or moral-dilemma positions along with negative sentiment co-occurring suggests that ethical risks are frequently perceived as an issue, and they are challenged within the community, which partly addresses the fears that GenAI is merely promoting unregulated abuse.

The creativity metrics represent student interaction as mostly informative but with significant creative pockets. Semantic and lexical diversity histograms have concentrated mid-range distributions, and it means that posts are more likely to discuss the same academic issues in a variant slightly varying in their constructions rather than talking about a completely novel issue. It fits the concept of ChatGPT as a study partner among the students when faced with repetitive tasks, such as essay plan, exam preparation, and concept explanation. The zero-shot creativity labels, where 78.1 percent of posts are found in the informative category and only 1.6 percent in the highly creative category, confirm the opinion that GenAI is not a generative, but a supportive, creative agent at the moment. However, the existence of original and funny labels suggests that other users explore more open and playful applications, which can be found in the literature that speaks of hybrid forms of human-AI creativity.

A tag topic subreddit co-occurrence network analysis indicates that discourse is structured around a few high-centrality hubs, in particular, Conversation, College, Work, International, and Social tags, and some prevalent topics. These hubs can be viewed as the places where local norms of the acceptable use of ChatGPT are constructed and disseminated. The observation of a few small communities indicates that various educational settings, including college admission, workplace upskilling, or foreign study, form somewhat different sentiment patterns, ethical orientation, and innovativeness.

On the whole, the findings confirm that the strength of ChatGPT utilization in student Reddit groups does not directly undermine creativity or ethics. Rather creativity seems to be directed to informational and problem-solving modes whereas ethical reflexivity is manifested in the discussions which directly negotiate the limits of reasonable use. Transformer-based classification, diversity measures, and network analysis are effective in terms of exploring these subtle patterns on a large scale, but qualitative inspection is necessary to gain insights into the model results and discover subtle ethical or creative indicators that automated pipelines might overlook.

8 Conclusion

This study analysed the use of ChatGPT among educational Reddit communities in learning and the role of creativity and ethical decision-making in the discourse of mass users. The study offers a comprehensive perspective of the embedding of generative AI in daily learning activities by using a mixed-methods design that integrated transformer-based sentiment and ethical stance classification, creativity metrics, and social network analysis on 177,485 posts.

The results indicate that the student talk about ChatGPT is mostly neutral or even positive in tone and mostly prosocial or ethically neutral in attitude, which means that GenAI is not primarily a source of polarised debate but a practical cognitive stimulus and together, these sentiment, ethical stance, creativity and network measures are the ways to measure creativity, ethical position and sentiment in Reddit discussions identified in the research question.. The analysis of creativity implies that the majority of interactions are informative and problem-oriented and fewer and significant cases of

original or humorous applications are used to demonstrate collaborative human-AI creativity. Network structures can be seen to display discourse as being structured around a few educational centres- College, Work, and Conversation- within which norms and practices around the use of ChatGPT are constantly negotiated.

TF-IDF using logistic regression as well as recurrent models like LSTMs were both thought of but not adopted as first line tools. These methods can deal with local context and frequency of words but fail to address the subtle sarcasm and long-range context of a Reddit conversation. They also necessitate data of task-specific training, which is not present with ethical stance and creativity labels in this area. Transformer models thus provided an improved trade-off between contextual and zero-shot flexibility and scalability on the 177,485-post dataset.

These findings can be used in academic discussions to offer empirical data that large language models can expand, but not necessarily harm, student agency when integrated into communities that actively discuss the ethical and responsible use of large language models. In terms of methodology, the research proves that transformer pipelines and network analysis can be joined to operationalise such complex constructs as creativity and ethics using unstructured social media data of scale.

The only limitation that was left is that this study did not compare the transformer pipeline to simpler baselines on this dataset. To compare sentiment and topic-level tasks to the models of TF-IDF, LSTM, or lightweight CNN, the future work can introduce the results of the work to these models. These experiments would measure the amount of extra explanatory power that the transformer models would provide over the traditional techniques.

Meanwhile, the constraints of possible model bias, the lack of subreddit metadata, and the use of zero-shot classification support the fact that one should interpret the results with great caution. These outcomes cannot be assumed as precise measures of actual creativity or moral thinking but as systematic indicators which refer to larger trends. These limitations can be overcome by future studies, and the framework can be applied to other platforms, user populations, and multimodal interactions to develop a more comprehensive picture of human-AI collaboration in education.

9 Future Work

A number of avenues come out to expand and enlarge this study. To begin with, subsequent research must include a time dimension to explicitly simulate the development of sentiment, ethical stance, and creativity over time with the introduction of new versions of ChatGPT and institutional policies. The existing analysis is rather cross-sectional; adding timestamps would allow the longitudinal models that would monitor the change in attitudes after significant events like policy announcements, exams, or high-profile cases of AI abuse.

Second, more research is required to lessen the reliance on zero-shot classification and enhance construct validity. Sentiment, ethical stance, and creativity Fine-tuning of transformer models and stricter assessment of classification performance in educational settings would be possible by gathering a manually annotated subset of posts. Multi-label ethics modelling that differentiates between such issues as plagiarism, privacy, fairness, and psychological dependence instead of putting them under one broad label can be supported by such gold-standard labels.

Third, the framework of creativity may be broadened in future studies by including textual diversity and zero-shot labels. As an example, higher-order creative processes could be better characterized by developing sequence-level measures of idea transformation, metaphor use or multi-turn patterns of interaction. The analysis of the texts and behavioural data, including editing paths, follow-up questions, or self-reported learning outcomes, would give a more complete idea about the impact of GenAI on the creative problem-solving of the students.

Fourth, the network analysis can be generalized through the reconstruction of full conversation trees and more detailed subreddit metadata. This would allow exploration of the ways in which certain communities: subject-specific subreddits or local groups would vary in their norms, and the ways influential users or posts would hasten or slow the adoption of particular ethical or creative practices. The addition of time-varying community detection would also be useful in explaining the solidification or disintegration of norms as the use of ChatGPT becomes more mature.

Lastly, there is opportunity to do translational work relating these findings to educational practice and policy. The information on locations of harmful or morally questionable applications may be used to construct specific digital-literacy interventions, recommendations to use ChatGPT in coursework, and resources to assist educators. Likewise, knowing in what areas truly creative applications are

flourishing may prompt educational architecture to purposefully tap into GenAI to collaboratively ideate, write, and reflect. The framework created beyond academia might be scaled to platform providers or regulators to track new patterns of risks in real time and create tools that encourage users to engage more responsibly and creatively with generative AI technologies.

References

- Buz T., Frost B., Genchev N., Schneider M., Kaffee L., Melo G. "Investigating Wit, Creativity, and Detectability of Large Language Models in Domain-Specific Writing Style Adaptation of Reddit's Showerthoughts" *Proceedings of the 2024 Annual Conference of the North American Chapter of the Association for Computational Linguistics*, 2024. Link: <https://aclanthology.org/2024.starsem-1.23.pdf>
- Chiarello F., Giordano V., Spada I., Barandoni S., Fantoni G. "Future applications of generative large language models: A data-driven case study on ChatGPT". *Technovation*, Volume 133 (May 2024). Link: <https://doi.org/10.1016/j.technovation.2024.103002>
- Chugh R., Turnbull D., Morshed A., Sabrina F., Azad S., Md Mamunur R., Kaisar S., Subramani S. "The Promise and pitfalls: a literature review of generative artificial intelligence as a learning assistant in ICT education" *Computer Applications in Engineering Education*, Volume 33, Issue 2 (March 2025). Link: <https://doi.org/10.1002/cae.70002>
- Desdevises J. "The paradox of creativity in generative AI: high performance, human-like bias, and limited differential evaluation" *Front. Psychol.* Volume 16 (2025). Link: <https://doi.org/10.3389/fpsyg.2025.1628486>
- Garrido A. P., Barra E. "Sentiment Analysis With Transformers Applied to Education: Systematic Review" *International Journal of Interactive Multimedia & Artificial Intelligence* 9.2 (February 2025). Link: <https://doi.org/10.9781/ijimai.2025.02.008>
- Gupta P., Ding B., Guan C., Ding D. "Generative AI: A systematic review using topic modelling techniques" *Data and Information Management*, Volume 8, Issue 2, June 2024. Link: <https://doi.org/10.1016/j.dim.2024.100066>
- Huang Y., Arora C., Huong W. C., Kanij T., Madulgalla A., Grundy J.. "Ethical Concerns of Generative AI and Mitigation Strategies: A Systematic Mapping Study" *January 2025. Available at SSRN (Social Science Research Network)*: <https://ssrn.com/abstract=5114708>
- Jackson V., Vasilescu B., Russo D., Ralph P., Izadi M., Prikladnicki R., D'angelo S., Inman S., Andrade A., A. van der Hoek. "The Impact of Generative AI on Creativity in Software Development: A Research Agenda." *Association for Computing Machinery Transactions on Software Engineering and Methodology*, Volume 34, Issue 5, Article No.: 133, Pages 1 - 28 (May 2025). Link: <https://doi.org/10.1145/3708523>
- Köckritz J., İlgen B., Cohrdes C., Hattab G. "Current applications and future directions in natural language processing for news media and mental health." *Sci Rep* 15, 32532 (2025). Link: <https://doi.org/10.1038/s41598-025-18413-z>
- Kowalkiewicz M. "The Economy of Algorithms: AI and the Rise of the Digital Minions". *The Economy of Algorithms: AI and the Rise of the Digital Minions*. Bristol, United Kingdom: Bristol University Press (2024). 224 pages. (Published on September 2025). Link: <https://procomm.ieee.org/book-reviews-ai-emerging-technologies/>
- Luo X., Ghosh S., Tilley J. L., Besada P., Wang J., Xiang Y.. "Shaping ChatGPT into my Digital Therapist": *A thematic analysis of social media discourse on using generative artificial intelligence for mental health* *Digital health* 11 (2025): 20552076251351088. Link: <https://doi.org/10.1177/20552076251351088>

Lv Z. “Generative artificial intelligence in the metaverse era” *Cognitive Robotics, Volume 3, July, 2023*. Link: <https://doi.org/10.1016/j.cogr.2023.06.001>

Masciari E., Umair A., Ullah H. “A Systematic Literature Review on AI-Based Recommendation Systems and Their Ethical Considerations” *IEEE Access (August 2024)*. Link: <https://ieeexplore.ieee.org/stamp/stamp.jsp?arnumber=10654261>

Mentzer K., Price J., Singh J. “Analyzing reddit discourse surrounding generative AI” *Issues in Information Systems, Volume 25, Issue 3 pp. 277-292 (2024)*. Link: https://doi.org/10.48009/3_iis_2024_122

Nedungadi P., Veena G., Tang K., Menon R. R. K., Raman R. “AI Techniques and Applications for Online Social Networks and Media: Insights from BERTopic Modeling” *IEEE Transactions on Network Science and Engineering, vol. 8, no. 4, 2022*. Link: <https://doi.org/10.1109/ACCESS.2025.3543795>

Pandy G., Pugazhenth V. J., Murugan A. “Generative AI: Transforming the Landscape of Creativity and Automation” *International Journal of Computer Applications (0975 - 8887) Volume 186 - No.63, January 2025*. Link: <https://www.ijcaonline.org/archives/volume186/number63/pandy-2025-ijca-924392.pdf>

Petratos P., Giannoula M. “Transformer-based AI for Sentiment Analysis in Marketing” *ACMLC '24: Proceedings of the 2024 6th Asia Conference on Machine Learning and Computing, Pages 177 - 185, March 2025*. Link: <https://doi.org/10.1145/3690771.3690795>

Ray P. P. “ChatGPT: A comprehensive review on background, applications, key challenges, bias, ethics, limitations and future scope” *Internet of Things and Cyber-Physical Systems Volume 3, 2023, Pages 121-154 (April, 2023)*. Link: <https://doi.org/10.1016/j.iotcps.2023.04.003>

Sengar S. S., Hassan A. B., Kumar S., Carroll F. “Generative Artificial Intelligence: A Systematic Review and Applications” *Multimedia Tools and Applications 84.21 (2025): 23661-23700. (May 2024)*. Link: <https://arxiv.org/html/2405.11029v1>

Talafidaryani M., Moro, S. “Public perception of ChatGPT on Reddit social media platform: Topic modeling and sentiment analysis study.” *Published on SSRN (Social Science Research Network) on February 2024*. Link: <https://doi.org/10.2139/ssrn.4716839>

Trigka M., Dritsas E. “The Evolution of Generative AI: Trends and Applications” *IEEE Access Volume: 13, pp. 98504-98529 (May 2025)*. Link: <https://doi.org/10.1109/ACCESS.2025.3574660>

Tunca S. “Algorithms of emotion: A hybrid NLP analysis of neurodivergent Reddit communities.” *Acta Psychologica, Volume 260, (October 2025)*. Link: <https://doi.org/10.1016/j.actpsy.2025.105519>

Xu Z., Fang Q., Huang Y., Xie M. “The public attitude towards ChatGPT on reddit: A study based on unsupervised learning from sentiment analysis and topic modeling” *Plos one, 19(5), e0302502 (May, 2024)*. Link: <https://doi.org/10.1371/journal.pone.0302502>