

Social Relationship and Content Analyses in an Online Class Community under Covid-19

Cindy Xindi Tong^{1 * †}, Rosanna Yuen-Yan Chan^{1 a †} and Choi Sen Ho^{2 b †}

¹ Department of Information Engineering, The Chinese University of Hong Kong

² Department of Computer Science and Engineering, The Chinese University of Hong Kong

* Corresponding author: Cindy Xindi Tong (Email: xtong308@gmail.com); ^a yychan@ie.cuhk.edu.hk; ^b luyou00001@gmail.com

[†] These authors contributed equally to this work

Abstract: Clarifying the blog's background and successfully grasping its main subject are essential components of any examination of the blogging community. There are numerous previous researches on how to evaluate blogs and comments to discover the connections and content in a social group. This work seeks to focus on blog comments based on students' relationships and conduct in a class community and to complement prior research techniques by emphasizing the importance of blog community analysis under Covid-19. It also provides an evaluation of the participation and performance results. We mainly focus on analyzing comments as the interaction among students and discuss the similarity in some highly performance students. With the help of social network analysis for online course student behavior, some tendency related to student behavior in this group has been shown in the experiment results by bias analysis, Pearson correlation, and Latent Dirichlet allocation, which serve as the data processing tools in this experiment.

Keywords: Social network, Online interaction; Social network analysis; Latent Dirichlet allocation; Pearson correlation.

1. Introduction

In 2020, a population lockdown was necessary in the majority of the world's nations to manage the coronavirus (COVID-19). The bulk of social and economic activities saw severe repercussions, and higher education institutions were forced to adjust to this unusual scenario by adjusting their teaching strategies and student evaluation procedures [1][2]. Both students and teachers must adjust to the sudden change when a course that was intended for face-to-face instruction is switched from an in-person to an online approach during the semester. Although it has been negative for education, it has provided the chance to examine any improvements to teaching and student evaluation methods that may be considered for inclusion in future courses once things return to normal [3][4].

Nowadays understanding how people interact with one another in web-based courses or lectures has become increasingly crucial given the growing importance of online education today. In a web-based graduate education course, we focus on the cognitive style profiles and language behaviors of students inside these communication networks. Group behavior in one community is usually analyzed from 9 major research areas, separately social cognition, attitudes, violence and aggression, prosocial behavior, prejudice and discrimination, social identity, group behavior, social influence, and interpersonal relationships. Many different new methods of interaction between humans and information have emerged from the combination of social actions and cyber-physical technology infrastructure. Under such circumstances, social groups on the Internet provide numerous research objects for understanding human behavior, including their behavior, opinions, preferences, personal, social, and cultural differences, and many other things. More research is needed to fully understand how individual characteristics affect student involvement, information-sharing patterns, and communication hierarchies in online

classrooms, as various scholars have noted. Using two complementary analytical methodologies [5].

In this paper, we characterize the individual behaviors connected to the communicative exchanges in asynchronous online discussion forums in an exploratory study to further this quest for understanding [5]. Instructors who are more conscious of how individual behaviors contribute to the creation of communication hierarchies and social roles in self-organizing groups might develop more effective strategies for structuring online activities that will promote higher student involvement [5].

1.1. Paper Organization

The rest of the paper is organized as follows: Section 2 provides an overview of the important ideas from the literature that underpins our theoretical frameworks and empirical methodology. Bias analysis, natural language processing using sentiment score and latent Dirichlet allocation, Graph theorem and NetworkX evaluation, and bias analysis are a few of them. We raise the research question of the current research and provide the methodology of our data processing strategy and data collecting in Section 3. The results of our experiment are displayed in Section 4. In Section 5, we will discuss the relations among our results. Finally, in Section 6, we will give the conclusion as well as the limits of this study and potential future research areas.

2. Preliminaries

2.1. Graph Theory and Network Analysis

Graph theory focuses on the study of graphs, modeling pairwise relations between objects through mathematical structures. A graph is constituted by vertices which could also be called nodes, and connected by edges which could also be called links or lines [6]. In the direction graph for a community, the nodes correspond to different users. In this way, users' comments sent and received can be displayed

through nodes: the indegree of one node means the comments the users receive, and the outdegree of one node signifies the number of comments the users send.

NetworkX is a Python package for complex network analysis, which is commonly used in graph creation and manipulation, containing network visualization and social network analysis. [8] On the one hand, network visualization aims to visualize the target network graph, including setting controls such as the color, shape, size, label, and any other components. On the other hand, social network analysis can help calculate overall network metrics, like density or modularity, and basic vertex metrics, such as degrees. Besides calculation, it is also able to complete group vertices by cluster or attributes [7].

2.2. Bias Analysis

Estimating potential biases' amount and direction as well as putting a number on their level of uncertainty is the main goal of bias analysis. There have been models to describe the direction and size of biases in recent years [8]. Systematic error or bias means that these deviations are not caused by accidental events. Instead, it shows a tendency in one research group that most of the participants would have, which is considered as the overlapping similarity for all the students in this group's behavior research (p.87).

The parameters for this experiment are set as follows: firstly, the j th comment length of the student with the class number i is ls_{ij} ; secondly, the length under blogs with the class number i is lb_{ij} ; thirdly, the number of comments that this student receives is NS_i and he or she sent is Nb_i .

Therefore, the formulas for the average result ls_i and lb_i for the student with class id i sending and receiving comments are

$$ls_i = \frac{1}{NS_i} \sum_{j=1}^{NS_i} ls_{ij} \quad (1)$$

$$lb_i = \frac{1}{Nb_i} \sum_{j=1}^{Nb_i} lb_{ij} \quad (2)$$

Then the bias of comment length for student i 's sending bs_i and receiving bb_i separately are

$$bs_i = \frac{1}{NS_i} \sum_{j=1}^{NS_i} (ls_{ij} - ls_i)^2 \quad (3)$$

$$bb_i = \frac{1}{Nb_i} \sum_{j=1}^{Nb_i} (lb_{ij} - lb_i)^2 \quad (4)$$

The bias is a criterion for deviation in a data group. If a smaller bias results in a control variable experiment, it means that the changing variable is more likely to be the main factor.

2.3. Natural Language Processing

Even though there is a lot of natural language text in the linked world and it contains a lot of knowledge, it is getting harder for humans to disperse it and find the knowledge or wisdom in it, especially given time constraints. The automated NLP is designed to perform this task as accurately and efficiently as a human would (for a limited of amount text). This chapter discusses the difficulties of NLP, the advancements made in the subject thus far, NLP applications, NLP components, and English grammar as it is needed by machines. Even though there is a lot of natural language text in the linked world and it contains a lot of knowledge, it is getting harder for humans to disperse it and find the knowledge or wisdom in it, especially given time constraints.

The automated NLP is designed to perform this task as accurately and efficiently as a human would (for a limited of amount text). This chapter discusses the difficulties of NLP, the advancements made in the subject thus far, NLP applications, NLP components, and English grammar as it is needed by machines [9].

Natural language processing (NLP) is a method of training the computer to process natural language input and generate the corresponding natural language output. It is a large field involving computer science, linguistics, artificial intelligence, human-computer interface, and some other areas (p.2).

The corresponding algorithm of NLP is to enable the computer to "understand" the content of documents which also includes the sight meaning of the language within them. Through this technique, it is capable of getting access to accurate information and topic collection involved in the target documents as well as document categorization and organization [10].

2.4. Sentiment Classification

The sentiment classification used in this research is Text Classification through a Dictionary-based Approach. It can be considered as polarity detection, which is aimed at estimating an opinion document and classifying writers' opinions into 3 different classes, separately positive, neutral, and negative [11].

We compute the sentiment scores according to the method introduced in [12]. Specifically, the computation process is based on a given dictionary, including positive and negative words, which respectively represent two poles of the detection. After obtaining the words set for the target document d , the calculation for the average 'sentiment score' needs to determine its polarity score by looking up the Dictionary for each word w . The score of w is -1 if it is a -ve word. The score of w is 0 if it is neither a +ve nor a -ve word. The score of w is 1 if it is a +ve word (Equation 5).

$$polarity(d) = \begin{cases} +ve & \text{if } d > 0 \\ -ve & \text{if } d < 0 \end{cases} \quad (5)$$

Then the normalized score is computed by dividing N , the total number of words in the target document (Equation 6).

$$s(d) = \frac{1}{N} \sum_{w \in d} score(w) \quad (6)$$

The normalized value is a score between -1 to +1 for objective comparisons in this group behavior research [11].

2.5. Latent Dirichlet Allocation

In natural language processing, the latent Dirichlet allocation (LDA) is a generative probabilistic model of a corpus to figure out the reasons why some parts of the content are similar by allowing observation sets to be explained by those unobserved groups. As one of the examples of a topic model, LDA functions on the principle that documents can be considered as random mixtures over latent topics, where the different topics are characterized by the distribution of words.

LDA assumes the following generative process for each document w in a corpus D . The first step is to choose $N \sim \text{Poisson}(\epsilon)$ and then choose $\theta \sim \text{Dir}(\alpha)$. And for each of the N words w_n , we need to choose a topic $z_n \sim \text{Multinomial}(\theta)$ and choose a word w_n from $p(w_n|z_n, \beta)$, a multinomial probability conditioned on the topic z_n [13].

This fundamental model contains several simplifying assumptions, some of which we eliminate in later parts. The dimensionality of the topic variable, z , and the Dirichlet distribution's dimensionality, k , are first assumed to be known

and fixed. Second, the word probabilities are parameterized by a $k \times V$ matrix, where $\beta_j = p(w_j = 1 | z_i = 1)$, which we treat as a constant that needs to be estimated. The Poisson assumption is not necessary for everything that comes after, and more accurate document length distributions are needed. Also, take note that N is independent of both z and any other variables that produce data. As a result, we will typically ignore its randomness in the development that follows because it is an auxiliary variable [14].

A k -dimensional Dirichlet random variable θ can take values in the $(k - 1)$ -simplex (a k -vector θ lies in the $(k - 1)$ -simplex if $\theta_i \geq 0, \sum_{i=0}^k \theta_i = 1$), and has the following probability density on this simplex:

$$P(\theta | \alpha) = \frac{\Gamma(\sum_{i=0}^k \alpha_i)}{\prod_{i=0}^k \Gamma(\alpha_i)} \theta_1^{\alpha_1-1} \dots \theta_k^{\alpha_k-1} \quad (7)$$

where the Gamma function is represented by $\Gamma(x)$ and the parameter is a k -vector with components $\alpha_i > 0$. The Dirichlet distribution is practical for the simplex because it belongs to the exponential family, has statistics with limited dimensions, and is conjugate to the multinomial distribution. The joint distribution of a topic mixture, a set of N topics, and a set of N words, given the parameters, is given by:

$$\prod_{n=1}^N p(\theta, z, w | \alpha, \beta) = p(\theta | \alpha) \prod_{n=1}^N p(z_n | \theta) p(w_n | z_n, \beta) \quad (8)$$

To capitalize on text-oriented intuitions, we refer to the latent multinomial variables in the LDA model as themes.

However, other than their usefulness in modeling probability distributions on sets of words, we make no epistemological claims about these latent variables where

$p(z_n | \theta)$ is simply θ_i for the unique i such that $z_n^i = 1$. Integrating over θ and summing over z , we obtain the marginal distribution of a document:

$$p(w | \alpha, \beta) = p(\theta | \alpha) \left(\prod_{n=1}^N \sum_{z_n} p(z_n | \theta) p(w_n | z_n, \beta) \right) d\theta \quad (9)$$

Finally, taking the product of the marginal probabilities of single documents, we obtain the probability of a corpus:

$$p(D | \alpha, \beta) = \prod_{d=1}^M \int p(\theta_d | \alpha) \left(\prod_{n=1}^{N_d} \sum_{z_{dn}} p(z_{dn} | \theta_d) p(w_{dn} | z_{dn}, \beta) \right) d\theta_d \quad (10)$$

The variables z_{dn} and w_{dn} are word-level variables and are sampled once for each word in each document. [14]

3. The Current Research

Our purpose for this research is to understand the social dynamics of one online-based blog community under the background of the epidemic. The following are the main research questions for this investigation:

Q1: How do social relationship affects individual participation in terms of social relationship and content popularity in terms of length and sentiment?

Q2: What is the relationship between the length and sentiment of our students?

Q3: Do influential participants emerge with similar content in their blogs?

The first 2 questions focus on the comment content for each individual's incoming and out coming comment text information. And question 3 needs the information in the first 2 questions' analyzing part, which is focused on the blog content summary.

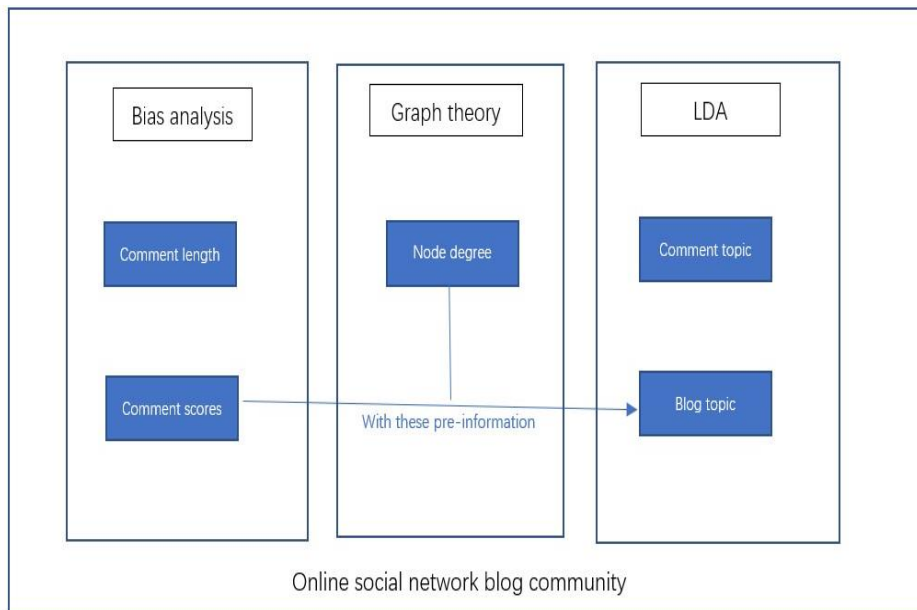


Fig. 1 The total research process in this essay

This study first analyzes the tendency of length and sentiment score and topic of their comments for different students and then digs deeper into the blogs for some of the students with some most outstanding characteristics in sentiment score, activity, and popularity. Following the preceding logic, we first calculate the average length and average sentiment scores for all the students received from and given to other students, then construct a general topic modeling for comments each student gives and receives. Additionally, the top performers in commenting length and

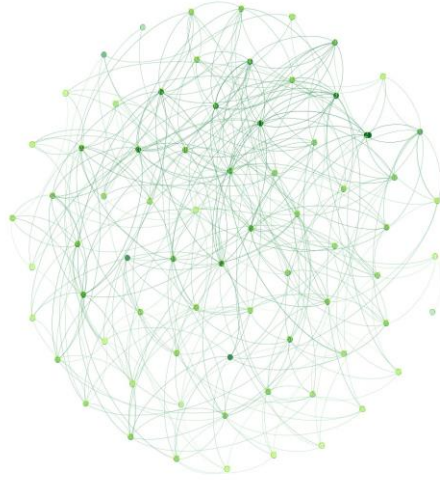
highly recommended are also collected as evidence. From the results of NetworkX for students with the highest indegree and outdegree and highest sentiment scores, some of the students' grading and discussion tendencies will be shown.

In Figure 1 showing the preceding logic, we first calculate the average length and average sentiment scores for all the students received from and given to other students, then construct a general topic modeling for comments each student gives and receives. Additionally, the top performers in commenting length and highly recommended are also

collected as evidence. From the results of NetworkX for students with the highest indegree and outdegree and highest sentiment scores, some of the students' grading and discussion tendencies will be shown.

4. Result

The data collected and processed mainly demonstrate two



points, firstly the overlapping points for all the students, and secondly the similarities and differences of some best performance students. And then we use the correlation table to show the relationship between the different parameters in this survey.

4.1. Social Network Analysis

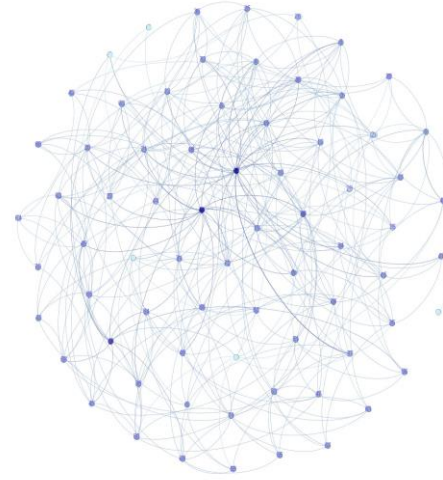


Fig. 2 Directed graph with greener node & **Fig. 3** Directed graph with darker node meaning higher outdegree meaning higher indegree

We first display the direct interaction status of the students with their group behavior. Figure 2 and Figure 3 show the indegree and outdegree of different students through the change in node color. In Figure 2, the higher indegree number of students is represented by the greener node while the yellow node means the student has a lower indegree. In Figure 3, the higher outdegree is presented by a darker node while the lighter one means the student has a lower outdegree. It is clear that almost all the students actively participate in class since there is seldom an isolated node or singular node. However, the range of the sent and received comment numbers varies among students. There exist some students

who sent or receive a large number of comments, which could be considered active or popular ones respectively.

4.2. Comment Length Analysis

In Figure 4 the average and bias for the length and sentiment scores for each student and each blog are first obtained showing respectively in the blue and yellow lines. The upper subplot in both Figure 3 is the average result and the latter one is the calculated bias. For the comment length, the bias for students writing which is the blue line in Figure 4 is always extremely smaller than the one under the same students' blog which is the orange line in the same figure.

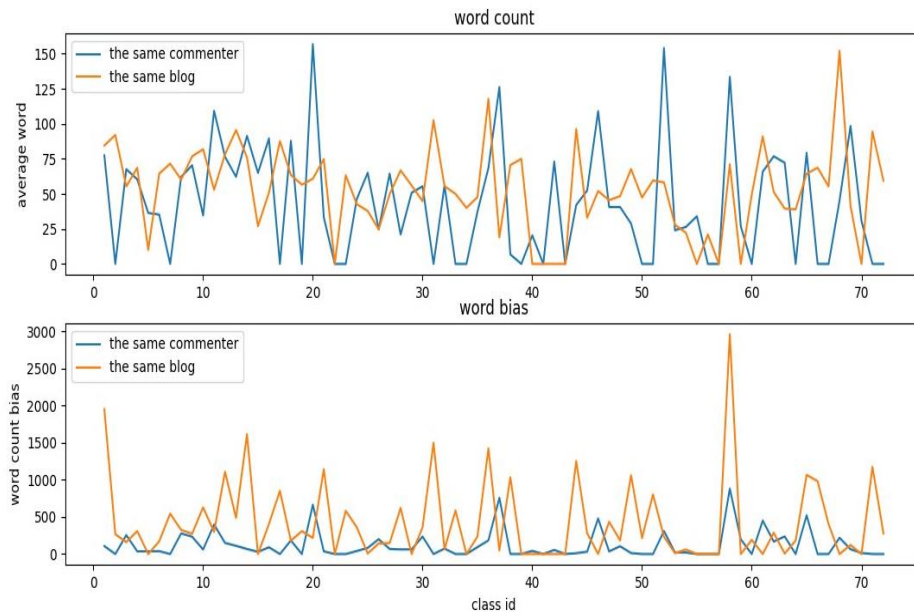


Fig. 4 Average comment word length and word bias for each student sent and received

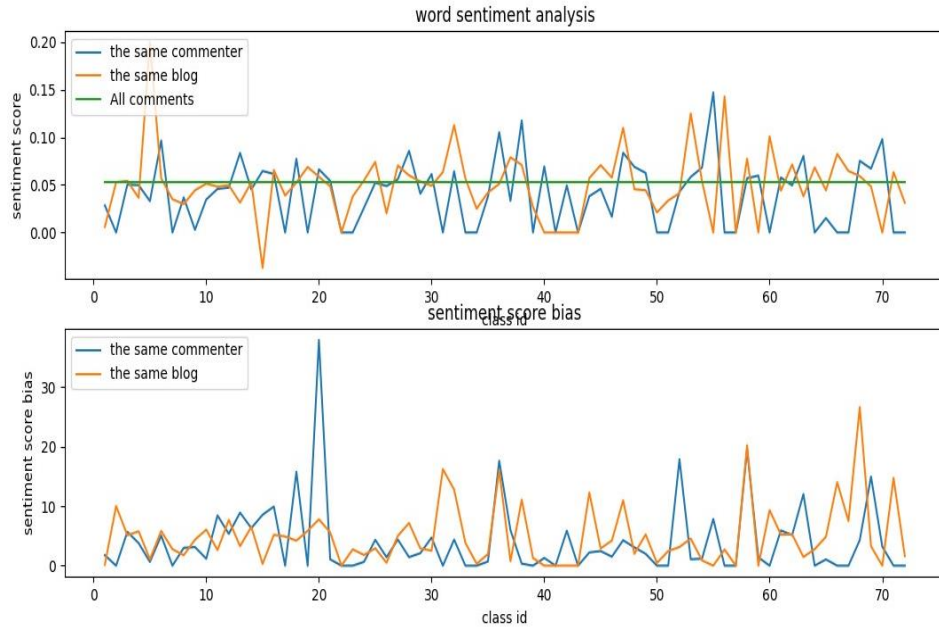


Fig. 5 Average comment sentiment score and score bias for each student sent and received

The case that this class group is one online group and students in this group could be considered as unfamiliar at the beginning of this course. Here we would define it as the interaction under different blogs. This indicates that the main factor for comment length may be the student's own choice, rather than the friendship under a blog. The meanings for the orange and blue lines correspond to the same research objects as in Figure 4.

4.3. Sentiment Analysis

For sentiment scores in Figure 5, the sentiment score

represented by the green line for all the blogs is larger than 0 while it is difficult to distinguish the main factor for the sentiment score since it does not show a stable gap between the bias. The sentiment score did not show a consistent tendency with the influence of friendship or personal choice. This may result from the factor for sentiment score tendency is not single, which may be the quality of the blogs.

4.4. Correlation among SNA Indexes, Comment Length, and Sentiment

Correlations													
		indegree	outdegree	closeness	betweenness	average length for sent comments	average length of received comments	bias of sent comments length	bias of received comments length	sentiment score of sent comments	sentiment score of received comments	bias of sent comments sentiment score	bias of received comments sentiment score
indegree	Pearson Correlation	1	.004	.790**	.679**	.299*	.111	.223	.094	.087	.127	.024	.407**
	Sig. (2-tailed)		.974	<.001	<.001	.011	.355	.059	.431	.465	.287	.844	<.001
	N	72	72	72	72	72	72	72	72	72	72	72	72
outdegree	Pearson Correlation	.004	1	.042	.576**	.166	-.021	-.098	-.016	.337**	.131	.154	-.043
	Sig. (2-tailed)	.974		.725	<.001	.162	.859	.412	.894	.004	.274	.196	.718
	N	72	72	72	72	72	72	72	72	72	72	72	72
closeness	Pearson Correlation	.790**	.042	1	.533**	.266*	.394**	.156	.147	.162	.372**	.085	.362**
	Sig. (2-tailed)	<.001	.725		<.001	.024	<.001	.191	.219	.174	.001	.477	.002
	N	72	72	72	72	72	72	72	72	72	72	72	72
betweenness	Pearson Correlation	.679**	.576**	.533**	1	.357**	.108	.115	.171	.221	.080	.113	.324**
	Sig. (2-tailed)	<.001	<.001	<.001		.002	.366	.337	.151	.063	.506	.347	.006
	N	72	72	72	72	72	72	72	72	72	72	72	72
average length for sent comments	Pearson Correlation	.299*	.166	.266*	.357**	1	.238*	.715**	.223	.278*	.044	-.028	.190
	Sig. (2-tailed)	.011	.162	.024	.002		.044	<.001	.060	.018	.714	.814	.110
	N	72	72	72	72	72	72	72	72	72	72	72	72
average length of received comments	Pearson Correlation	.111	-.021	.394**	.108	.238*	1	.283*	.728**	.029	-.045	.208	-.035
	Sig. (2-tailed)	.355	.859	<.001	.366	.044		.016	<.001	.812	.706	.079	.770
	N	72	72	72	72	72	72	72	72	72	72	72	72
bias of sent comments length	Pearson Correlation	.223	-.098	.156	.115	.715**	.283*	1	.226	-.047	-.255*	-.038	.245*
	Sig. (2-tailed)	.059	.412	.191	.337	<.001	.016		.056	.692	.030	.751	.038
	N	72	72	72	72	72	72	72	72	72	72	72	72
bias of received comments length	Pearson Correlation	.094	-.016	.147	.171	.233	.728**	.226	1	-.082	-.249*	.159	-.011
	Sig. (2-tailed)	.431	.894	.219	.151	.060	<.001	.056		.496	.035	.183	.927
	N	72	72	72	72	72	72	72	72	72	72	72	72
sentiment score of sent comments	Pearson Correlation	.087	.337**	.162	.221	.278*	.029	-.047	-.082	1	.229	.426**	.034
	Sig. (2-tailed)	.465	.004	.174	.063	.018	.812	.692	.496		.053	<.001	.780
	N	72	72	72	72	72	72	72	72	72	72	72	72
sentiment score of received comments	Pearson Correlation	.127	.131	.372**	.080	.044	-.045	-.255*	-.249*	.229	1	.023	.015
	Sig. (2-tailed)	.287	.274	.001	.506	.714	.706	.030	.035	.053		.847	.899
	N	72	72	72	72	72	72	72	72	72	72	72	72
bias of sent comments sentiment score	Pearson Correlation	.024	.154	.085	.113	-.028	.208	-.038	.159	.426**	.023	1	.015
	Sig. (2-tailed)	.844	.196	.477	.347	.814	.079	.751	.183	<.001	.847		.902
	N	72	72	72	72	72	72	72	72	72	72	72	72
bias of received comments sentiment score	Pearson Correlation	.407**	-.043	.362**	.324**	.190	-.035	.245*	-.011	.034	.015	.015	1
	Sig. (2-tailed)	<.001	.718	.002	.006	.110	.770	.038	.927	.780	.899	.902	
	N	72	72	72	72	72	72	72	72	72	72	72	72

** Correlation is significant at the 0.01 level (2-tailed).
* Correlation is significant at the 0.05 level (2-tailed).

Fig. 6 Pearson correlation between degree-related variables and experiment results

How the relationship of different variables with some evaluation like Pearson correlation, sig, and the total number of students N. There is a total of 72 students in the online

course. Correlation results signed with * means there exist a statistically significant relationship between the 2 corresponding factors.

5. Discussion

5.1. Dynamics among Friendship, Comment Length, and Sentiment

Here we define friendship as the input and output of each student, which could be also considered as the number of sending number and receiving the number Social Relationship and Content Analyses

of comments for each student. Friendship in this survey could be reflected in the degree, closeness, and betweenness of students.

A more direct result could be got from the correlation table. Despite the conception decided by related parameters like closeness and betweenness with indegree and outdegree. The average of received comments length is moderately correlated with indegree and betweenness and it is considered highly positively correlated with closeness. The average of sent comments length is considered highly positively correlated with closeness, which means those students who show high closeness, also could be considered as the more active students in this survey tended to send longer comments. The average sentiment score of sent comments is highly positively correlated with outdegree and the average sentiment score of received comments is highly positively correlated with closeness, showing that the more active student were tending to receive high comments. The bias of sent comments does not have a relation with degree-related parameters while the bias of received comments shows a highly positive relation with indegree, closeness, and betweenness. This may be due to the more students comment on one blog, the more possible different options on the content.

The sent and receive comments status will affect each other in the corresponding part. There are relations between the average length of sent and received message comments and the bias of these comments' length. The average of sent comments length is moderately positively correlated with the average of received comments length and is highly positively correlated with the bias of received comment length. The bias of sent comment length is highly positively correlated with the bias of received comment length. From these relation data, we can infer that students who wrote longer comments will also have a high possibility to receive longer comments.

To summarize these data, those students be more popular or have more friends would like to give respond to others actively and have a higher possibility to receive longer comments and higher sentiment feedback.

5.2. Relation between Comment Length and Sentiment

There are also relations between average sentiment sent and received message comments and the bias of these scores. The average of the sent comment sentiment score is highly positively correlated with the bias of the sent comments sentiment score. And this means that those who like to give higher sentiment scores would also tend to give different evaluations of others' work.

The comment length and comment scores also affect each other. The average of sent comments length is moderately positively correlated with the average of sent comments sentiment score. The bias of sent comment length is highly positively correlated with the bias of sent comments sentiment score. The average received sentiment score is moderately negatively correlated with the bias of sent and received comment length. This means that those who tend to

write longer comments to others would also prefer to give others higher evaluation, those who did not fix their writing length would also tend to not provide stable feedback, and those who got higher evaluation would like to give their comments with less fixed comments length.

5.3. Characteristic of Popular Blogs

For one thing, the overlapping topic in these blogs is answering the question because of the course blog writing requirements. This shows that almost all the students in one group show a similar analyzing methodology for comments. Additionally, the receiving topics also contain the question as a high-frequency topic while there are still some different topics ranging in all the different chapters.

For another, some differences are also shown in the topic for the most highly rated students between the most active ones and the most popular ones. For highly graded students, their blog shows the preference for the calculation process according to the high-frequency mathematics characters. For those students who are the most active or most popular, they show more preference for concept explanation because there are more concept terms related to this course in their topics' distribution. The reason for the difference is that the calculation process leads to definite answers while concept explanation provides more discussion space, as the former remains fixed and the latter keeps flexible.

Though there are some differences in the topic distribution for the students above, they share the same topic and information, which is the central concept of the class blog community.

Besides some basic information collected through statistics, such as the most popular one and the average sentiment score of comments, there are some other content-related findings for this class blog community. Firstly, the topic of all the blogs is directly related to the central point of this course, following speculation. Secondly, for the preference in one blog community, students still hold the tendency to provide feedback according to their ideas when compared with peer pressure in the same blog.

6. Conclusion

During the research process, we found that the general possible factor for the tendency of word length is a personal choice. And both word length and sentiment score depend on the popularity of the student. The parameter of comment length and sentiment score will also affect each other. The more active and more serious the student is in his comments writing, the higher possibility he will receive longer comments and higher evaluation on this coursework. The data provided by LDA also reveal that the highest-rated students and the most active or popular students hold different focuses or preferences on formula calculation and concept explanation.

From our research evidence, the students still maintained social activities through online class community during Covid-19. More importantly, the sentiment was positive in general. Due to the limitations of time and analysis Social Relationship and Content Analyses

material, some preset questions have not been answered thoroughly. For example, the tendency of the length of students' comments and attitudes is still affected by several other reasons such as students' motion at that time and their workload, so it maybe is considered as one multi factors problem. The possible causes of these problems will be

determined in the future with more time and advanced research tools.

Author Contributions

Cindy Xindi Tong and Rosanna Yuen-Yan CHAN wrote the main manuscript text. Choi Sen HO prepared data for figure 1-2. All authors reviewed the manuscript

Funding

No funding.

Data Availability Statement

Data derived from public domain resources.

Conflicts of Interest

The authors declare no conflict of interest.

Appendix A

LDA result of selected students' blog

Topic modeling result for the most outstanding student in 3 areas Highest scores:

[(0, '0.134**x" + 0.044**figure" + 0.036**p" + 0.036**log" + 0.031**like"), (1, '0.080**h" + 0.058**state" + 0.039**subject" + 0.032**id" + 0.027**question"), (2, '0.037**agent" + 0.031**understanding" + 0.031**brain" + 0.029**q" + 0.026**information")] Most active:

[(0, '0.051**information" + 0.026**value" + 0.022**entropy" + 0.021**twin" + 0.017**log"), (1, '0.144**h" + 0.028**p" + 0.027**different" + 0.019**message" + 0.014**energy"), (2, '0.148**x" + 0.034**mutual" + 0.027**step" + 0.022**bit" + 0.016**e")] Most popular:

[(0, '0.083**information" + 0.031**agent" + 0.030**free" + 0.025**theory" + 0.023**blog"), (1, '0.034**energy" + 0.033**x" + 0.026**understanding" + 0.025**understand" + 0.018**try"), (2, '0.024**mutual" + 0.021**h" + 0.017**variable" + 0.014**random" + 0.011**word")]

References

- [1] Haushofer, J., Metcalf, C.J.E.: Which interventions work best in a pandemic? *Science* 368(6495), 1063–1065 (2020).
- [2] Walker, P.G., Whittaker, C., Watson, O.J., Baguelin, M., Winskill, P., Hamlet, A., Djafaara, B.A., Cucunub'a, Z., Olivera Mesa, D., Green, W., et al.: The impact of covid-19 and strategies for mitigation and suppression in low-and middle-income countries. *Science* 369(6502), 413–422 (2020).
- [3] Santos, H.: Covid-19 lockdown effects on student grades of a university engineering course: A psychometric study. *IEEE Transactions on Education* (2021).
- [4] Drury, J., Carter, H., Cocking, C., Ntontis, E., Tekin Guven, S., Aml'ot, R.: Facilitating collective psychosocial resilience in the public in emergencies: Twelve recommendations based on the social identity approach. *Frontiers in public health* 7, 141 (2019).
- [5] Vercellone-Smith, P., Jablow, K., Friedel, C.: Characterizing communication networks in a web-based classroom: Cognitive styles and linguistic behavior of self-organizing groups in online discussions. *Computers & Education* 59(2), 222–235 (2012).
- [6] Biggs, N., Lloyd, E.K., Wilson, R.J.: *Graph Theory, 1736-1936*. Oxford University Press (1986).
- [7] Hagberg, A., Swart, P., S Chult, D.: *Exploring network structure, dynamics, and function using NetworkX* (2008).
- [8] James, L.R., Demaree, R.G., Wolf, G.: Estimating within-group interrater reliability with and without response bias. *Journal of applied psychology* 69(1), 85 (1984).
- [9] Hinton, A.: *Understanding context: Environment, language, and information architecture.* O'Reilly Media, Inc." (2014).
- [10] Chowdhary, K.: *Natural language processing. Fundamentals of artificial intelligence*, 603–649 (2020).
- [11] Mikolov, T., Chen, K., Corrado, G., Dean, J.: Efficient estimation of word representations in vector space. *arXiv preprint arXiv:1301.3781* (2013).
- [12] Kiritchenko, S., Zhu, X., Mohammad, S.M.: Sentiment analysis of short informal texts. *Journal of Artificial Intelligence Research* 50, 723–762 (2014).
- [13] Blei, D.M., Ng, A.Y., Jordan, M.I.: Latent dirichlet allocation. *Journal of machine Learning research* 3(Jan), 993–1022 (2003)
- [14] Benesty, J., Chen, J., Huang, Y., Cohen, I.: *Pearson Correlation Coefficient*, Springer (2009).