

Sentiment Analysis Related to Law No. 6 Of 2023 on the Employment Cluster Using the Bidirectional Long Short-Term Memory Algorithm

Anugra M^{1*}, Munawar²

¹²Program Studi Teknik Informatika, Fakultas Ilmu Komputer, Universitas Esa Unggul

DOI:

<https://doi.org/10.53697/jkomitek.v4i2.2051>

*Correspondence: Anugra M
anugerahmatippanna@gmail.com

Received: 20-10-2024

Accepted: 20-11-2024

Published: 21-12-2024



Copyright: © 2024 by the authors. Submitted for open access publication under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

Abstract: The enactment of UU Cipta Kerja has triggered varied public responses, particularly concerning employment provisions like fixed-term employment (PKWT), the legalization of outsourcing, unfair severance pay, and ease of layoffs. Social media has become a primary platform for the public to share opinions on issues within the law's employment cluster. This study employs sentiment analysis using the Bidirectional Long Short-Term Memory (Bi-LSTM) algorithm to understand public sentiment about UU Cipta Kerja and sentiment within its content. Bi-LSTM is chosen for its ability to capture temporal relationships and context in long texts, which aids in handling complex sentiment classification. The findings indicate varied public perceptions: neutral sentiment dominates issues like "PKWT" and "Minimum Wage" on Twitter (X), reflecting uncertainty. Positive sentiment appears around "Outsourcing" and "Minimum Wage" provisions, indicating perceived flexibility. Conversely, negative sentiment dominates issues like "Layoffs" and "Severance," on both social media and in UU Cipta Kerja content, signaling concerns over worker rights. The Bi-LSTM model achieved 70.15% accuracy for the Twitter dataset and 83.22% for the law's content dataset.

Keywords: Sentiment Analysis, Cipta Kerja, Bi-LSTM

Introduction

At the beginning of 2023, the Indonesian House of Representatives (DPR) passed Law No. 6 of 2023 on Job Creation to attract investors, aiming to stimulate job creation and deliver economic benefits to society (Matompo & Izziyana, 2020). The Omnibus Law approach was employed in drafting this law to achieve regulatory effectiveness, efficiency, and harmonization (Situngkir, 2022). The government considered this step crucial due to Indonesia's high unemployment rate, which reached eight million people (Kurniawan, 2020).

However, the enactment of the Job Creation Law faced significant criticism, particularly regarding the Employment Cluster, which was perceived to undermine labor protection guarantees (Iswaningsih et al., 2021). According to a survey conducted by Kompas Research and Development (Sakti, 2023) of 512 respondents, 49% believed the law

was not pro-labor, 19% felt it made layoffs easier for companies, 11% claimed it directly affected their jobs, 5% argued it influenced maximum contract duration limits, and 16% believed it pressured workers in their jobs. Key issues raised concerning the enactment of Law No. 6 of 2023 include (Wicaksono, 2023):

- a. Ambiguity in regulations regarding the completion timeframe for work under Fixed-Term Employment Agreements (PKWT).
- b. Simplified access for companies to outsource workers for their core business activities.
- c. Uncertainty in minimum wage policies for workers.
- d. Eased regulations enabling companies to terminate employment (layoffs).
- e. Perceived unfair compensation for workers affected by layoffs.

These issues led to protests by various groups, including workers and students, as well as judicial reviews filed by the Labor Party with the Constitutional Court. The enactment of the Job Creation Law by the DPR also drew diverse responses and became a hot topic among users on the social media platform Twitter X. This research examines the contents of Law No. 6 of 2023 and public tweets related to the Employment Cluster of the Job Creation Law as a source of opinions and sentiments on the law's enactment using sentiment analysis methods. Sentiment analysis involves processing unstructured text to predict subjective expressions categorized as positive, negative, or neutral (Romadhan et al., 2023). This study focuses on utilizing deep learning methods, specifically the implementation of Word2Vec for word embedding and Bidirectional LSTM (BI-LSTM) for data training. The results are expected to provide reliable predictions in identifying sentiments related to the contents of Law No. 6 of 2023 using the BI-LSTM model.

Several previous studies on sentiment analysis have been conducted. For instance, (Sarimole & Ihsan, 2023) analyzed sentiment on Twitter concerning the Job Creation Law using Naïve Bayes and Support Vector Machine algorithms. In that study, data labeling was automated with TextBlob. The Naïve Bayes algorithm achieved an accuracy of 75.43%, precision of 84.80%, recall of 65.75%, and an F1-score of 66.03%. The Support Vector Machine algorithm yielded an accuracy of 81.31%, precision of 83.14%, recall of 75.73%, and an F1-score of 77.49%. Additionally, (Wijaya et al., 2021) conducted a sentiment analysis of public opinion on the Job Creation Law on Twitter using the Naïve Bayes algorithm, achieving an accuracy of 89.9%, precision of 90%, recall of 89.9%, and an F1-score of 89.9% with 80% training data and 20% testing data. (Dauni et al., 2022) analyzed Twitter user sentiment toward the Omnibus Law Bill using Convolutional Neural Networks (CNN). In that study, training data of 2,538 records yielded an average accuracy of 78.6% with 10 epochs. Using 20 epochs achieved the highest accuracy of 90% with 90% training data, and an average accuracy of 86.8% over five tests. (Sandryan et al., 2021) performed a sentiment analysis on Twitter regarding the Job Creation Law using the Backpropagation algorithm and Term Frequency-Inverse Document Frequency (TF-IDF). The best classification accuracy was 98%, obtained with a learning rate of 0.2, 10 hidden nodes, and 1,500 epochs. Precision was 98%, recall was 92.4%, and the F1-measure was 95.1%. Based on previous research, this study employs BI-LSTM for sentiment analysis related to Law No. 6 of 2023,

aiming to better understand public sentiments and opinions regarding the enactment of the Job Creation Law’s Employment Cluster.

Methodology

This research combines Text Mining and Sentiment Analysis methods, with the object of study being the content of Law No. 6 of 2023 on the Employment Cluster. The study involves several stages, including Data Collection, Data Preprocessing, Word Embedding, Data Labeling, Wordcloud Visualization, Classification using Deep Learning, and Model Evaluation. The research employs Python programming language with the assistance of Google Colaboratory.

Data Collection

The collection of content from the Job Creation Law was sourced from the JDIH website of the Audit Board of the Republic of Indonesia, which was then converted into Excel file format. This process resulted in 468 pieces of data on the Employment Cluster, and after removing duplicates, 466 final data entries were obtained. Tweets were collected using scraping techniques with the "Tweet Harvest" program to gather tweets based on specific keywords within a particular timeframe. Keywords such as "PKWT," "Outsourcing," "PHK," "Pesangon," "Upah," and "Cipta Kerja" were used to collect tweets related to the Employment Cluster of the Job Creation Law from January 1, 2023, to May 30, 2023. After duplicate removal, a total of 10,588 tweets were gathered.

Text Preprocessing

Text preprocessing involves transforming and cleaning data into a structured format. This stage removes punctuation, symbols, numbers, and irrelevant words, and converts words to their base forms. Preprocessed data are then ready for word embedding models to represent words as vectors for sentiment classification. The preprocessing steps include:

- 1. *Cleansing*: Removes non-alphabetic characters such as symbols and punctuation to minimize noise.
- 2. *Case Folding*: Converts all characters to lowercase for text consistency.
- 3. *Stopword Removal*: Eliminates common words like conjunctions and pronouns to focus on more informative terms.
- 4. *Stemming*: Strips affixes from words, returning them to their root form.
- 5. *Tokenizing*: Splits sentences into individual words for separate analysis.

Table 1. Results of Text Preprocessing

Content	Preprocessed Text
Employers must pay compensation for each foreign worker employed.	['employers', 'must', 'pay', 'compensation', 'for', 'foreign', 'workers']
Employment relationships arise due to a work agreement between employers and workers.	['employment', 'relationship', 'arises', 'agreement', 'employers', 'workers']

Word Embedding

Word embedding is an important technique in Natural Language Processing (NLP) that aims to represent words in the form of low-dimensional vectors (Aspri et al., 2020). This representation allows computers to understand and manipulate text more effectively. Word2Vec is one of the popular algorithms used to generate word embeddings, which aims to transform words into vector representations in a low-dimensional space (Patil et al., 2023). This algorithm utilizes an unsupervised learning approach to learn the meanings of words from the context in which they appear (Chen & Sokolova, 2021).

Data Labeling

Data labeling is the process of assigning sentiment labels to each piece of data. In this study, the data labeling process involves categorizing the data into three types of sentiment: "Positive," "Negative," and "Neutral." The data labeling in this study uses the VADER Lexicon method. VADER is a method that works by assigning a sentiment score to each word based on a word dictionary within VADER. When analyzing a sentence in a tweet, each word in the sentence is assigned a sentiment score as defined in the VADER lexicon. These scores are then summed up to get a combined compound score for the sentence. If the final combined score is greater than 0.1, the tweet is classified as positive sentiment; if the result is less than -0.05, the tweet is classified as negative sentiment; and if the result falls between -0.05 and 0.1, the tweet is classified as neutral sentiment (Hutto & Gilbert, 2014). Based on the data labeling results using VADER, the dataset for the Omnibus Law (UU Cipta Kerja) contains 72 data points with negative sentiment, 157 data points with neutral sentiment, and 238 data points with positive sentiment. For the Twitter (X) dataset, there are 3408 data points with negative sentiment, 1944 data points with neutral sentiment, and 5205 data points with positive sentiment. The sentiment labeling results using the VADER Lexicon are shown in Table 2.

Table 2. Data Labeling

Content	Label
Employers must pay compensation for each foreign worker employed.	Neutral
The obligation to pay compensation as referred to in paragraph (1) does not apply to government agencies, foreign representatives, international bodies, social institutions, religious organizations, and certain educational institution positions.	Neutral
Employment relationships arise due to a work agreement between employers and workers.	Positive

Wordcloud Visualization

After the data labeling process, a word cloud visualization is performed to display the most frequently occurring words in the dataset. Word cloud visualization helps in understanding the information within the dataset by representing the most frequent words. The word cloud provides an appealing visual representation of word frequency, where words that appear more often are displayed in larger sizes

Data Splitting

Oversampling is a technique used to address the class imbalance problem in a dataset. Class imbalance occurs when one class in the dataset has significantly more examples than

the other classes. Oversampling works by increasing the number of examples in the minority class until it matches the majority class, either by duplicating existing data or generating new data based on the existing data (Nurhopipah & Magnolia, 2023). One popular oversampling method is Synthetic Minority Over-sampling Technique (SMOTE), which increases the number of examples in the minority class by creating synthetic examples (Nugroho & Rilvani, 2023). After the data labeling and oversampling processes, the dataset is split into two parts: the training data and the test data. The training data is used to train the deep learning model, which is used for sentiment prediction in sentiment analysis. The test data is used to evaluate the performance of the trained model in predicting sentiment on new data. The data is split with 80% allocated for training and 20% for testing.

Classification Model Implementation

Bidirectional Long Short-Term Memory (BiLSTM) is an extension of LSTM that processes sequential data in two directions: from front to back and vice versa. This allows BiLSTM to capture context from both directions, providing a richer representation of sequential data (Onan, 2022). BiLSTM is highly useful in natural language processing, including sentiment analysis, due to its ability to consider information from both previous and subsequent contexts simultaneously. BiLSTM is often used in applications that require a deeper understanding of the order of words in a text.

In addition to its ability to capture context from both directions, BiLSTM is also effective in addressing the vanishing gradient problem that often occurs in conventional neural networks when processing long sequences. This is because LSTM, including BiLSTM, uses a gating mechanism that adaptively controls the flow of information and gradients through the network. By considering context from both directions, BiLSTM enables the model to make more accurate and precise decisions in determining sentiment polarity. This advantage makes BiLSTM particularly effective for tasks such as named entity recognition, machine translation, and other tasks that require a complex understanding of context.

BiLSTM can also work well with various types of word embeddings, such as Word2Vec or GloVe, which enrich the feature representation of text. Therefore, BiLSTM has become a popular choice in many NLP applications that require in-depth contextual understanding.

Model Evaluation (Confusion Matrix)

The Confusion Matrix is an important evaluation tool in analyzing the performance of classification models, including in the context of sentiment analysis (Kolo & Supatman, 2024). The Confusion Matrix consists of four parts: True Positive (TP), False Positive (FP), True Negative (TN), and False Negative (FN). TP represents the number of cases correctly predicted as positive, while FP is the number of cases incorrectly predicted as positive. TN represents the number of cases correctly predicted as negative, while FN is the number of cases incorrectly predicted as negative (Arifqi et al., 2024). From the Confusion Matrix, we can calculate evaluation metrics such as accuracy, precision, recall, and F1-score, which provide insights into the model's performance in classifying sentiment (Citra R et al., 2024).

Table 3. Confusion Matrix

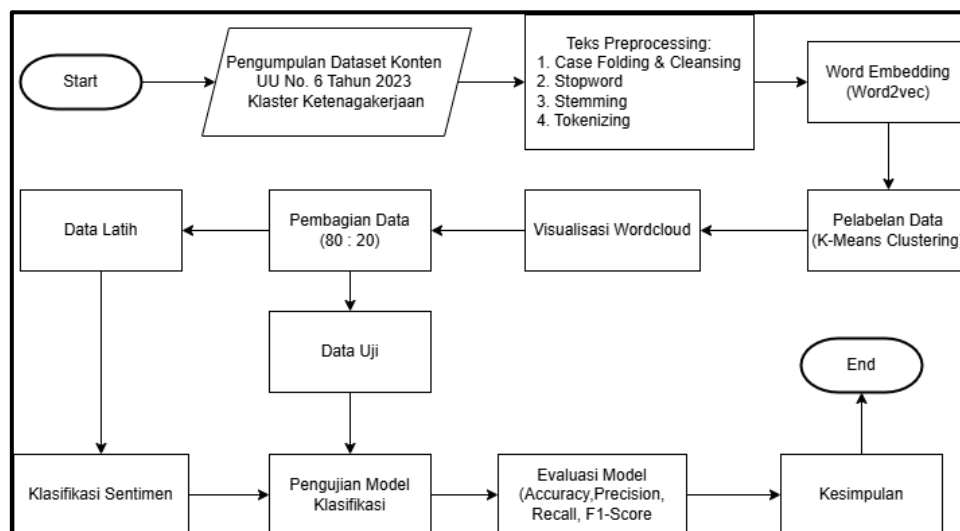
True Class	Predicted Class
	Positive
Positive	True Positive (TP)
Negative	False Positive (FP)
Neutral	False Positive (FP)

$$Accuracy = \frac{TP + TN + TNt}{TP + TN + TNt + FP + FN + FNT}$$

$$Precision = \frac{TP}{TP + FP}$$

$$Recall = \frac{TP}{TP + FN + FNT}$$

$$F1 \text{ Score} = 2 \times \frac{Precision \times Recall}{Precision + Recall}$$

**Figure 1.** Research Flow

Result and Discussion

Classification Model Results

The model used for the UU Cipta Kerja content data is Bidirectional LSTM with 1 layer consisting of 32 units. This layer functions to capture sequential information from both directions (forward and backward) in the data. After the Bi-LSTM layer, there is a Dropout layer with a rate of 50%, used for regularization to prevent overfitting by reducing the dependence between neurons. The model then has a Dense layer with 32 units that uses the ReLU activation function to introduce non-linearity into the model. This Dense layer is also equipped with L2 regularization ($\lambda=0.4$) to prevent weight values from becoming too large, which also helps to prevent overfitting. The last layer is a Dense output layer using the softmax activation function. This layer has 3 units, corresponding to the number of unique

classes in the target data (for example, for sentiment classification with 3 classes: positive, negative, and neutral).

The model is then compiled using the `sparse_categorical_crossentropy` loss function, the Adam optimizer with a learning rate of 0.001, a batch size of 8, and 80 epochs for iteration. The model used for Twitter(X) data is also a Bidirectional LSTM with 1 layer consisting of 32 units. This layer captures sequential patterns from both directions (forward and backward) in the data. After the LSTM layer, there is a Dropout layer with a rate of 50% for regularization, which helps prevent overfitting by reducing dependencies between neurons. The model continues with a Dense layer with 16 units and uses the ReLU activation function to introduce non-linearity into the model. This Dense layer is also equipped with L2 regularization ($\lambda=0.01$) to control weight values, preventing them from becoming too large, which also helps to prevent overfitting. The final layer is a Dense output layer using the softmax activation function. This layer has the number of units corresponding to the number of unique classes in the target data. The model is then compiled using the `sparse_categorical_crossentropy` loss function, the Adam optimizer with a learning rate of 0.0007, a batch size of 32, and 80 epochs for iteration. The confusion matrix results from testing the Bi-LSTM model for the UU Cipta Kerja Employment Cluster content dataset and Twitter X data can be seen in Figure 2.

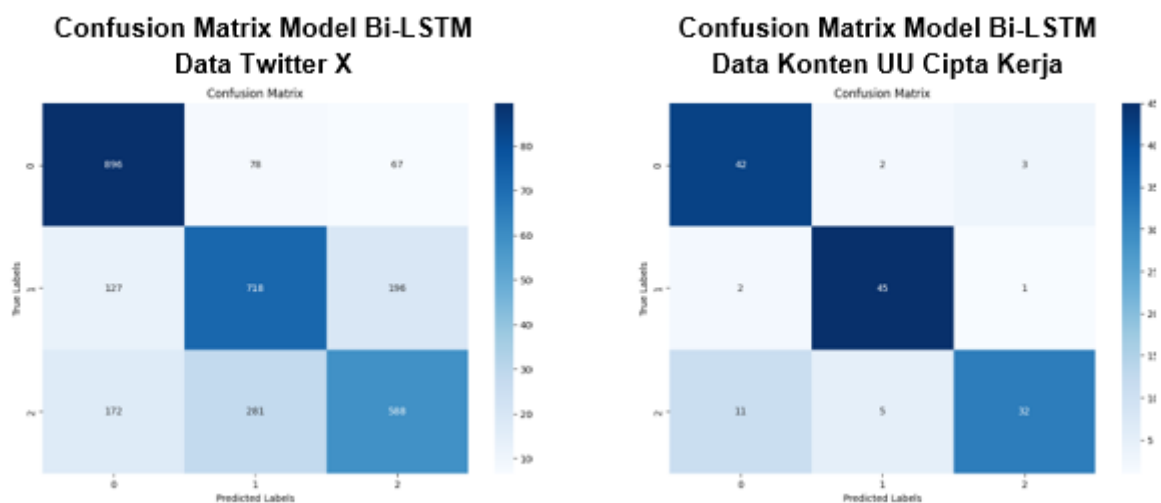


Figure 2. Bi-LSTM Confusion Matrix

Table 4. Precision, Recall, F1-Score of UU Cipta Kerja Content Data

Class	Performance	Precision	Recall	F1-Score
Negative (Label 1)	0.7636	0.8936	0.8235	
Neutral (Label 0)	0.8654	0.9375	0.900	
Positive (Label 2)	0.8889	0.6667	0.7619	
Accuracy	0.8322			

The performance of the Bi-LSTM model for the UU Cipta Kerja Employment Cluster content dataset based on the confusion matrix visualization results is presented in Table 4. Based on the test results, the final accuracy for model testing is 83.22%, which indicates that

the model has a high accuracy in classifying sentiments correctly. Precision is used to describe the accuracy of the data requested with the results provided by the model, calculated through the ratio of correct predictions compared to the total predicted results. The test results show a precision of 76.36% for the negative label, 86.54% for the neutral label, and 88.89% for the positive label. This means that the proportion of correctly predicted labels from the total predictions is very good for all three classes, especially for the negative label class. Recall describes the model's success in retrieving information from the test data, where recall is obtained by calculating the ratio of correct predictions compared to the total correct data. The test results show that the success rate of the model in retrieving information is 89.36% for the negative label, 93.75% for the neutral label, and 66.67% for the positive label. This indicates that the model's success in retrieving information is quite good for all three classes, especially for the negative and neutral labels.

Table 5. Precision, Recall, F1-Score of Twitter (X) Data

Class	Performance	Precision	Recall	F1-Score
Negative (Label 1)	0.6667	0.6897	0.6780	
Neutral (Label 0)	0.7498	0.8607	0.8014	
Positive (Label 2)	0.6910	0.5648	0.6216	
Accuracy	0.7015			

The performance of the Bi-LSTM model for Twitter (X) based on the confusion matrix visualization results is presented in Table 5. Based on the test results, the final accuracy for model testing is 70.15%, which indicates that the model has a good accuracy in classifying sentiments correctly. Precision is used to describe the accuracy of the data requested with the results provided by the model, calculated through the ratio of correct predictions compared to the total predicted results. The test results show a precision of 66.67% for the negative label, 74.98% for the neutral label, and 69.10% for the positive label. This means that the proportion of correctly predicted labels from the total predictions is very good for all three classes, especially for the positive label class. Recall describes the model's success in retrieving information from the test data, where recall is obtained by calculating the ratio of correct predictions compared to the total correct data. The test results show that the model's success rate in retrieving information is 68.97% for the negative label, 86.07% for the neutral label, and 56.48% for the positive label. This indicates that the model's success in retrieving information is very good for the negative and neutral labels. The loss and accuracy curves during the model training process can be seen in Figure 3.

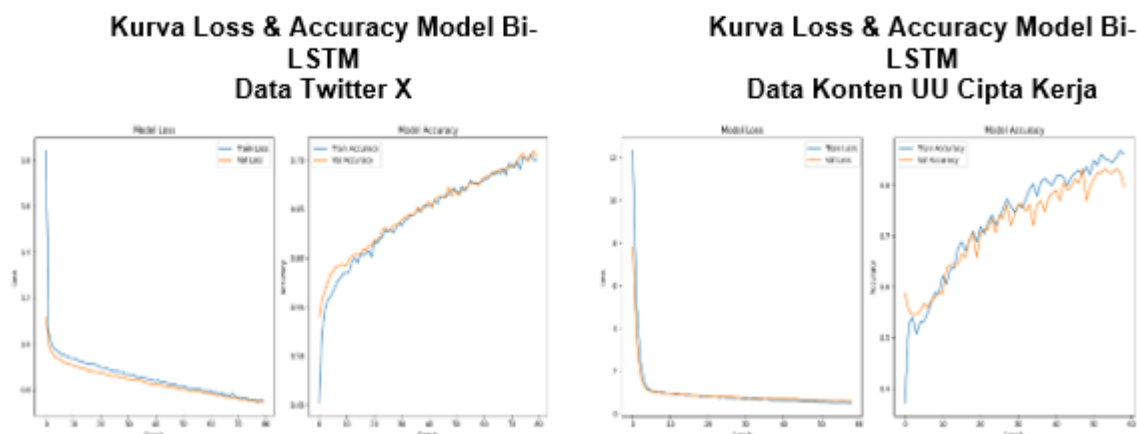


Figure 3. Bi-LSTM Loss and Accuracy Curves

Based on the curves above, the loss graph shows that both train loss and val loss continue to decrease as epochs increase. This indicates that the model is improving in minimizing prediction errors during training and validation. The accuracy graph shows that both train accuracy and val accuracy increase as epochs increase. This indicates that the model is improving in predicting data during training and validation. The curves above suggest a consistent trend, with no significant fluctuations, indicating stable training.

Visualization

Based on the data labeling results performed in the previous step, here are the Wordcloud visualization results based on sentiment labels. Figure 4 shows the Wordcloud for Twitter and UU Cipta Kerja content data with a dominant negative sentiment related to employment. On Twitter, words like "phk," "pecat," and "pesangon" reflect worker dissatisfaction with job insecurity and the often unfavorable contract system. Meanwhile, the UU Cipta Kerja content highlights issues like labor relations, employment regulations, and legal sanctions with words like "pekerja," "serikat," and "tindak pidana." Both sources indicate strong dissatisfaction from workers regarding working conditions and policies, with criticism of the protection of workers' rights being deemed insufficient. Figure 5 shows the Wordcloud for Twitter and UU Cipta Kerja content data with neutral sentiment, featuring words frequently appearing in descriptive and informative contexts without strong sentiment. On Twitter, words like "outsourcing," "pkwt," and "gaji" reflect discussions on various forms of employment, finances, and job status without emphasizing any particular preference or judgment. Meanwhile, UU Cipta Kerja content focuses on terms like "kerja," "pekerja," and "perusahaan," which are related to employment regulations and workers' rights in an objective manner. Both sources demonstrate neutral discussions focusing on informative aspects of employment, without emphasizing positive or negative sentiment. Example tweets and neutral UU Cipta Kerja content reflect labor regulations and the relationship between workers and employers.

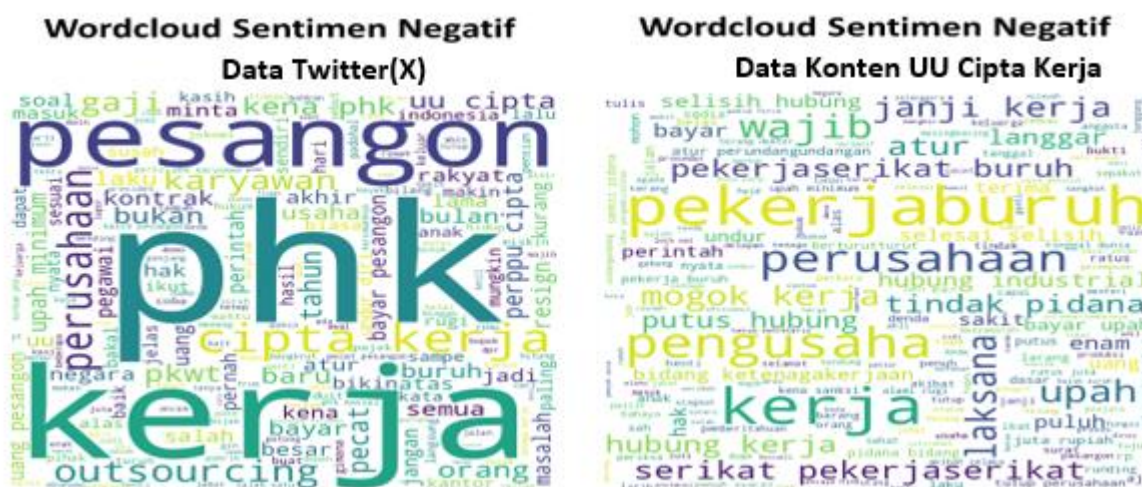


Figure 4. Negative Label Wordcloud



Figure 5. Neutral Label Wordcloud

Figure 5 represents the Wordcloud for Twitter and UU Cipta Kerja content data with a positive sentiment, displaying words that often appear in favorable and supportive contexts. On Twitter, words like "kerja," "ciptakerja," and "gaji" reflect a positive view of job opportunities, new labor regulations, and financial benefits for workers. Terms like "outsourcing" and "kontrak" also show a positive perspective on the legality and benefits of various forms of employment. On the other hand, UU Cipta Kerja content emphasizes words like "kerja," "serikat pekerja," and "janji," reflecting commitment to protecting workers' rights and employers' obligations. Words like "lembaga" and "struktur" indicate efforts to improve labor regulations. Both sources emphasize the benefits, support, and improvements expected from the UU Cipta Kerja, with examples like workers' rights and legal protections.



The word cloud from Twitter data and the content of the Omnibus Law (UU Cipta Kerja) with positive sentiment displays words that frequently appear in contexts that are favorable and supportive. On Twitter, words like "kerja" (work), "ciptakerja" (Cipta Kerja), and "gaji" (salary) reflect a positive view of job opportunities, new labor regulations, and financial benefits for workers. Terms like "outsourcing" and "kontrak" (contract) also indicate a positive perspective on the legality and advantages of various types of employment. On the other hand, the content of the Omnibus Law highlights words like "kerja" (work), "serikat pekerja" (labor union), and "janji" (promise), which reflect a commitment to worker rights protection and employer obligations. Words like "lembaga" (institution) and "struktur" (structure) indicate efforts to improve labor regulation. Both sources emphasize the benefits, support, and improvements expected from the Omnibus Law, with examples such as job training programs aimed at enhancing competence and productivity.

Conclusion

The conclusion of this study indicates that sentiment analysis using the Bidirectional Long Short-Term Memory (Bi-LSTM) algorithm is highly effective in understanding public sentiment regarding labor provisions in Law No. 6 of 2023 on Employment Cluster. From the analyzed data, which includes 488 content data of the law and 10,558 Twitter data, the Bi-LSTM model achieved an accuracy of 70.15% for Twitter data (X) and 83.22% for the content data of the Omnibus Law on Employment. The sentiment analysis results on various labor-related topics from Twitter data and the content of the Omnibus Law show significant variation in views. Fixed-Term Work Agreements (PKWT) tend to be neutral in sentiment on Twitter but positive in the law content, indicating that concerns about unilateral termination and compensation rights are well-regulated in the law. Outsourcing has a positive sentiment in both sources, with the Omnibus Law regulating provisions effectively and providing protection for workers' rights. Minimum Wage shows a positive sentiment

on Twitter and neutral sentiment in the Omnibus Law content, with the law offering flexibility in determining wages based on economic conditions. Termination of Employment (PHK) is dominated by negative sentiment on Twitter but positive in expert analysis of the Omnibus Law, emphasizing that the regulations do not permit unilateral termination without clear reasons. Severance pay generally has negative sentiment on Twitter and in the content of the Omnibus Law but positive in expert views, with the law clearly regulating compensation, though concerns about the amount of compensation remain.

Suggestions for future research include considering the use of other models, such as transformer-based models, to compare performance and effectiveness in sentiment analysis. Models like BERT (Bidirectional Encoder Representations from Transformers) and GPT (Generative Pre-trained Transformer) may provide better accuracy and deeper understanding of text context. Moreover, further research could expand the data volume and diversify data sources to gain a more comprehensive view of public sentiment. The implementation of more advanced models and improvements in data collection methods will enhance the validity and reliability of research findings, offering a more significant contribution to the field of sentiment analysis in public policy.

References

- Arifqi, T., Suarna, N., & Prihartono, W. (2024). Penggunaan Naive Bayes Dalam Menganalisis Sentimen Ulasan Aplikasi Mcdonald's Di Indonesia. *JATI (Jurnal Mahasiswa Teknik Informatika)*, 8(2), Article 2. <https://doi.org/10.36040/jati.v8i2.8740>
- Aspri, M., Tsagakatakis, G., & Tsakalides, P. (2020). Distributed Training and Inference of Deep Learning Models for Multi-Modal Land Cover Classification. *Remote Sensing*, 12(17), Article 17. <https://doi.org/10.3390/rs12172670>
- Chen, Q., & Sokolova, M. (2021). Specialists, Scientists, and Sentiments: Word2Vec and Doc2Vec in Analysis of Scientific and Medical Texts. *SN Computer Science*, 2(5), 414. <https://doi.org/10.1007/s42979-021-00807-1>
- Citra R, F., Indriyani, F., & Rahadjeng, I. R. (2024). Klasifikasi Tumor Otak Berbasis Magnetic Resonance Imaging Menggunakan Algoritma Convolutional Neural Network. *Digital Transformation Technology*, 3(2), 918–924. <https://doi.org/10.47709/digitech.v3i2.3469>
- Dauni, P., Ramdhani, M. A., Suryapratama, D. R., Jumadi, Zulfikar, W. B., & Fuadi, R. S. (2022). Twitter User Sentiment Analysis For RUU Omnibus Law Using Convolutional Neural Network. *2022 IEEE 8th International Conference on Computing, Engineering and Design (ICCED)*, 1–6. <https://doi.org/10.1109/ICCED56140.2022.10010404>
- Hutto, C., & Gilbert, E. (2014). VADER: A Parsimonious Rule-Based Model for Sentiment Analysis of Social Media Text. *Proceedings of the International AAAI Conference on Web and Social Media*, 8(1), Article 1. <https://doi.org/10.1609/icwsm.v8i1.14550>
- Iswaningsih, M. L., Budiarta, I. N. P., & Ujianti, N. M. P. (2021). Perlindungan Hukum Terhadap Tenaga Kerja Lokal dalam Undang-undang Nomor 11 Tahun 2020 tentang

- Omnibus Law Cipta Kerja. *Jurnal Preferensi Hukum*, 2(3), Article 3. <https://doi.org/10.22225/jph.2.3.3986.478-484>
- Kolo, S. Y., & Supatman, S. (2024). Analisis Sentimen Terhadap Opini Masyarakat Terkait Perubahan Cuaca Di Indonesia Menggunakan Algoritma Support Vector Machine. *JATI (Jurnal Mahasiswa Teknik Informatika)*, 8(2), Article 2. <https://doi.org/10.36040/jati.v8i2.8988>
- Kurniawan, F. (2020). Problematika Pembentukan RUU Cipta Kerja Dengan Konsep Omnibus Law. *Jurnal Panorama Hukum*, 5(1), Article 1. <https://doi.org/10.21067/jph.v5i1.4437>
- Matompo, O. S., & Izziyana, W. vivid. (2020). Konsep Omnibus Law Dan Permasalahan Ruu Cipta Kerja. *RECHTSTAAT NIEUW*, 5(1), Article 1. <https://ejournal.unsa.ac.id/index.php/rechtstaat-niew/article/view/506>
- Nugroho, A., & Rilvani, E. (2023). Penerapan Metode Oversampling SMOTE Pada Algoritma Random Forest Untuk Prediksi Kebangkrutan Perusahaan. *Techno. Com*, 22(1). <https://search.ebscohost.com/login.aspx?direct=true&profile=ehost&scope=site&authType=crawler&jrnl=14122693&AN=164068345&h=DscQNqI8urjYQhwBgRJn01b8IKXnXur1VhCm9uw8PIUdj2ucwMG4rYznnY%2B5Lb4CrIEiBjD7P6FPQGBB623Q%3D%3D&crl=c>
- Nurhopipah, A., & Magnolia, C. (2023). Perbandingan Metode Resampling Pada Imbalanced Dataset Untuk Klasifikasi Komentar Program MBKM. *Jurnal Publikasi Ilmu Komputer Dan Multimedia*, 2(1), Article 1. <https://doi.org/10.55606/jupikom.v2i1.862>
- Onan, A. (2022). Bidirectional convolutional recurrent neural network architecture with group-wise enhancement mechanism for text sentiment classification. *Journal of King Saud University - Computer and Information Sciences*, 34(5), 2098–2117. <https://doi.org/10.1016/j.jksuci.2022.02.025>
- Patil, R., Boit, S., Gudivada, V., & Nandigam, J. (2023). A Survey of Text Representation and Embedding Techniques in NLP. *IEEE Access*, 11, 36120–36146. IEEE Access. <https://doi.org/10.1109/ACCESS.2023.3266377>
- Romadhan, A. N., Utami, E., & Hartanto, A. D. (2023). Analisis Sentimen Opini Publik Menggunakan Metode BiLSTM Pada Media Sosial Twitter. *Seminar Nasional Teknologi Informasi Dan Matematika (SEMIOTIKA)*, 2(1), 22–33. <https://journal.itk.ac.id/index.php/semiotika/article/download/997/524>
- Sakti, R. E. (2023, January 16). *Perppu Cipta Kerja, Kekhawatiran dan Harapan*. [kompas.id](https://www.kompas.id/baca/riset/2023/01/16/perppu-cipta-kerja-kekhawatiran-dan-harapan). <https://www.kompas.id/baca/riset/2023/01/16/perppu-cipta-kerja-kekhawatiran-dan-harapan>
- Sandryan, M. K., Rahayudi, B., & Ratnawati, D. E. (2021). Analisis Sentimen Pada Media Sosial Twitter Terhadap Undang-Undang Cipta Kerja Menggunakan Algoritma Backpropagation dan Term Frequency-Inverse Document Frequency. *Jurnal Pengembangan Teknologi Informasi dan Ilmu Komputer*, 5(2), Article 2.

- Sarimole, F. M., & Ihsan, A. N. (2023). Analisis Sentimen Twitter Terhadap Uu Cipta Kerja Dengan Menggunakan Algoritma Naïve Bayes Dan Support Vector Machine. *INTECOMS: Journal of Information Technology and Computer Science*, 6(2), 822–829. <https://doi.org/10.31539/intecom.v6i2.7671>
- Situngkir, R. (2022). Urgensi Penerapan Omnibus law Untuk Menyelesaikan Permasalahan Pembentukan Regulasi Di Indonesia. *Iuris Studia: Jurnal Kajian Hukum*, 3(1), 1–8. <https://doi.org/10.55357/is.v3i1.193>
- Wicaksono, A. (2023). *Massa Buruh Geruduk MK-Istana Besok, Tuntut UU Ciptaker Dicabut*. <https://www.cnnindonesia.com/nasional/20230808175145-20-983530/massa-buruh-geruduk-mk-istana-besok-tuntut-uu-ciptaker-dicabut>
- Wijaya, T. N., Indriati, R., & Muzaki, M. N. (2021). Analisis Sentimen Opini Publik Tentang Undang-Undang Cipta Kerja Pada Twitter. *Jambura Journal of Electrical and Electronics Engineering*, 3(2), Article 2. <https://doi.org/10.37905/jjee.v3i2.10885>