



Hyperparameter Tuning Pada Model Stance Detection Menggunakan GridSearchCV

Arif Hamied Nababan^{1*}, Mia Yovanca Hutagalung²

^{1,2}Teknik Informatika, Universitas Prima Indonesia, Medan, Indonesia

Email Korespondensi : arifhamiednababan@unprimdn.ac.id

Abstrak– Tujuan dari penelitian ini adalah membangun model stance detection terhadap Undang-Undang Cipta Kerja yang memiliki performa lebih baik daripada model sebelumnya. Performa model *stance detection* ini penting untuk ditingkatkan agar dapat dilakukan penyerapan aspirasi masyarakat yang lebih baik terhadap Undang-Undang Cipta Kerja yang terus menjadi kontroversi besar di Indonesia mulai dari tahun 2019 sampai di tahun 2023 ini. Dalam penelitian ini digunakan metode *Hyperparameter Tuning* yang diimplementasikan menggunakan kerangka kerja CRISP-DM. Hasil yang diperoleh adalah sebuah model *stance detection* menggunakan algoritma *Support Vector Machine* yang memiliki nilai micro f1-score sebesar 78,4%. *Hyperparameter* model tersebut adalah C bernilai 10, parameter max_features dari proses ekstraksi fitur menggunakan metode TF-IDF sebesar 10000, dan menggunakan fitur unigram. Disimpulkan bahwa model tersebut memiliki performa lebih baik dengan nilai micro f1-score yang lebih besar 6,6% daripada model *stance detection* pada penelitian sebelumnya.

Kata Kunci: Data Science, Stance detection, UU Cipta Kerja, Hyperparameter tuning, Gridsearch

Abstract– The purpose of this research is to build a stance detection model for the Job Creation Law that has better performance than the previous model. The performance of this stance detection model is important to be improved so that better absorption of public aspirations can be carried out on the Job Creation Law which continues to be a big controversy in Indonesia from 2019 to 2023. In this research, the Hyperparameter Tuning method is used which is implemented using the CRISP-DM framework. The result obtained is a stance detection model using the Support Vector Machine algorithm that has a micro f1-score value of 78.4%. The hyperparameters of the model are C of 10, max_features parameter from the feature extraction process using TF-IDF method of 10000, and using unigram features. It is concluded that the model has better performance with a micro f1-score value that is 6.6% greater than the stance detection model in previous research.

Keywords: Data Science, Stance detection, Job Creation Bill, Hyperparameter tuning, Gridsearch

I. PENDAHULUAN

Undang-Undang (UU) Cipta Kerja disusun dengan menggunakan metode omnibus law, yaitu metode pemuatan sebuah undang-undang yang berisi berbagai macam permasalahan yang berbeda. UU Cipta kerja terdiri dari 11 klaster, 15 bab, dan 174 pasal serta memengaruhi 1.203 pasal dari 79 Undang-Undang. Beberapa ketentuan dari UU Ketenagakerjaan, UU perlindungan dan Pengelolaan Lingkungan Hidup, UU Penataan Ruang, dan penghapusan UU Administrasi Pemerintahan. Dalam konsiderannya, RUU Cipta Kerja disusun antara lain untuk menyerap bayak tenaga kerja, kemudahan berusaha, dan penyesuaian berbagai aspek pengaturan yang berkaitan dengan kemudahan berusaha [1].

Undang-Undang Cipta kerja adalah salah satu undang-undang yang menjadi prioritas dalam prolegnas periode 2020-2024 [2]. UU Cipta Kerja merupakan salah satu UU yang sangat kontroversial di Indonesia. Mulai dari proses penyusunan, isi, dan pengesahan dari UU Cipta kerja ini mendapat banyak sekali penolakan dari masyarakat [3]. Penolakan ini berujung pada terjadinya banyak demonstrasi di berbagai wilayah di Indonesia dan munculnya berbagai ekspresi penolakan maupun dukungan masyarakat di media sosial.

Walaupun menghadapi banyak penolakan, pada 30 Desember 2020 Undang-Undang Cipta kerja disahkan oleh DPR. Namun pengesahan tersebut tidak mengurangi penolakan masyarakat yang menganggap UU Cipta Kerja merugikan secara materiil dan memiliki cacat hukum dalam proses pembuatannya. Setelah melalui uji materi oleh Mahkamah Konstitusi pada 27 November 2021 memutuskan bahwa UU Cipta Kerja inkonstitusi bersyarat atau bertentangan dengan UUD 1945 [4]. Sampai pada tahun 2023, Perpu pengganti UU Cipta Kerja disahkan, namun tetap saja masyarakat membuat permohonan melakukan uji materi terhadap Perpu ini.

Melihat fenomena tersebut penelitian terdahulu membangun sebuah model *stance detection* terhadap UU Cipta Kerja [5]. Model ini ditujukan untuk dapat digunakan dalam memperluas dan mempercepat penyerapan aspirasi masyarakat terhadap UU Cipta Kerja. Model yang dibangun dapat dengan baik memprediksi keberpihakan pengguna media sosial Twitter berdasarkan *tweet* yang diposting selama periode penyusunan sampai pengesahan UU Cipta Kerja. Namun, melihat kontroversi dari UU ini masih berlangsung sampai tahun 2023, maka peningkatan performa model *stance detection* tersebut akan memberikan dampak yang lebih besar terhadap penyerapan aspirasi masyarakat.

Untuk itu pada penelitian ini akan dibangun sebuah model *stance detection* yang dapat mendeteksi



keberpihakan pengguna media sosial Twitter dengan performa lebih baik dari penelitian awal.

Stance detection adalah masalah klasifikasi di mana pendirian dari pengarang sebuah teks dilihat dalam bentuk beberapa kategori menggunakan set label: {Mendukung, Menentang, Tidak keduanya}[6]. Pada penelitian sebelumnya [5], tugas *stance detection* yang dibangun termasuk menggunakan *content* dari teks sebagai fitur yang akan dipelajari. Selain itu penelitian tersebut termasuk pada kelompok *target-specific stance detection*, yaitu tugas *stance detection* yang teks (T) atau pengguna (U) merupakan masukan utama untuk memprediksi *stance* terhadap target tunggal yang sudah ditentukan sebelumnya (G). Jika terdapat lebih dari satu target, maka model *stance detection* yang dibangun harus terpisah untuk masing-masing target [7].

Pada penelitian sebelumnya *Stance detection* sudah banyak diterapkan untuk menentukan pendirian atau keberpihakan dari penulis sebuah teks. Contohnya pada [8] sebuah model *stance detection* berhasil dibangun untuk menentukan penulis dari artikel blog mendukung seorang tokoh politik untuk memenangkan pemilihan umum. Mereka menggunakan kombinasi beberapa metode ekstraksi fitur dan memperoleh model terbaik dengan macro-average f1-score sebesar 63,54% dengan menggunakan algoritma Support Vector Machine (SVM) sebagai *classifier* dengan kombinasi fitur Word2vec dan Unigram.

Pada penelitian “*Twitter stance detection towards Job Creation Bill*” model *stance detection* dapat dibangun dengan baik [5]. Penelitian tersebut mencoba membangun sembilan buah model dengan menggunakan kombinasi Algoritma Multinomial Naïve Bayes, Logistic Regression, dan Support Vector Machine dengan menggunakan metode fitur ekstraksi TF-IDF yang menghasilkan vektor Unigram, bigram, dan Unigram+bigram. Selain itu parameter max_feature yang digunakan sebesar 5000 dan 10000 fitur. Model terbaik pada penelitian tersebut berhasil mencapai F1-Score sebesar 71.8%. Yaitu model dengan menggunakan Algoritma Logistic Regression dengan fitur Unigram+bigram dan menggunakan 10000 pada parameter max_featurenya.

Pada penelitian [9] mengimplementasi *Stance detection* menggunakan dua buah dataset, yaitu dataset SemEval 2016 yang berisi *tweet* mengenai lima topik tertentu. Kemudian mereka juga menggunakan dataset Multi Perspective Consumer Health Query yang berisi teks formal mengenai lima buah kueri yang berhubungan dengan kesehatan. Selain menggunakan beberapa metode standar, mereka juga menggunakan model pre-trained BERT dan menggunakan *Hyperparameter tuning* untuk untuk *Stance detection*. Hasil penelitian mereka menunjukkan bahwa penggunaan BERT yang menggunakan *hyperparameter tuning* memberikan performa lebih baik dari metode lainnya.

Penggunaan *Hyperparameter tuning* juga digunakan oleh [10] dalam mengklasifikasi ujaran kebencian di Twitter menggunakan algoritma Support Vector Machine. Model terbaik yang dihasilkan adalah model yang

menggunakan kernel sigmoid, parameter regularization sebesar 10, dan gamma bernilai 0.1. Akurasi yang diperoleh sebesar 74.88%

Metode *Hyperparameter tuning* pada *Classifier* Naïve bayes dan SVM juga dilakukan oleh [11]. Mereka melakukan analisis sentimen terhadap saham perusahaan-perusahaan kesehatan di Malaysia. Dalam implementasinya metode grid search digunakan dalam proses menentukan hyperparameter terbaik. Hasil yang diperoleh adalah mereka menemukan bahwa hyperparameter yang penting untuk *classifier* Naïve Bayes adalah alpha dan *fit_prior*. Sedangkan untuk *classifier* SVM, hyperparameter C, kernel, dan gamma menjadi hyperparameter terpenting. SVM memberikan model yang memiliki performa lebih baik setelah dilakukan *hyperparameter tuning*, yaitu dengan akurasi 85,65%.

Hyperparameter tuning pada *classifier* Logistic Regression juga dilakukan oleh [12]. Penelitian ini menganalisis sentimen masyarakat terhadap kebijakan vaksinasi covid-19 pada media sosial Twitter. Model klasifikasi yang dihasilkan memiliki f1-score sebesar 60% setelah dilakukan Slicing, dan *Hyperparameter tuning*.

Penggunaan metode grid search dalam implementasi hyperparameter tuning juga sudah banyak dilakukan, Salah satunya adalah [13]. Penelitian tersebut melakukan prediksi terhadap sentimen pelanggan dengan menggunakan algoritma Random Forest. Penggunaan grid search dan hyperparameter tuning memberikan hasil yang sangat menjanjikan. Akurasi model yang diperoleh tanpa hyperparameter tuning sebesar 84.53, tetapi ketika dilakukan hyper parameter tuning menggunakan grid search akurasi model meningkat hingga 90.02% [13].

Dengan demikian penelitian ini akan menggunakan hasil penelitian [5] sebagai tolok ukur model *stance detection*. Penelitian [8] digunakan sebagai dasar penggunaan *stance detection* dalam ranah politik dan kebijakan pemerintahan. Penelitian ini akan menggunakan metode *Hyperparameter tuning* berdasarkan pada penelitian [9]. Akan digunakan juga Algoritma SVM, Multinomial Naïve Bayes, dan Logistic Regression dalam metode Hyperparameter Tuning yang sudah teruji di penelitian [10], [11], dan [12]. Selain itu akan digunakan juga metode *grid search* dalam implementasi Hyperparameter tuning berdasarkan penelitian [13].

II. METODE PENELITIAN

Metode yang digunakan pada penelitian ini menggunakan CRISP-DM. Kerangka kerja tersebut merupakan salah satu yang paling populer dalam mengimplementasikan algoritma dan teknik-teknik khusus dalam mengekstrak pengetahuan dari data. CRISP-DM terdiri dari enam fase yang dilakukan secara sequensial namun bersifat iteratif di setiap prosesnya. Keenam fase tersebut adalah: *Business Understanding*, *Data Understanding*, *Data preparation*, *Modeling*, *Evaluation*, dan *Deployment* [14].



1. Business Understanding

Pada fase ini dilakukan pemahaman tentang latar belakang, permasalahan, dan tujuan dari penelitian. Hal-hal tersebut kemudian dikonversi menjadi sebuah permasalahan *data science* atau *machine learning*. Kemudian didesain sebuah rencana mengenai kebutuhan data, metode yang digunakan, dan metrik luaran dari penelitian. Selain itu pada fase ini juga dipertimbangkan *feasibility* penelitian yang akan dilakukan.

2. Data Understanding

Dataset yang digunakan berasal dari penelitian [5]. Dataset tersebut berisi teks *tweet* dari media sosial Twitter. Dataset dianotasi oleh sembilan orang annotator. Data awal berjumlah 1.114.773 namun dilakukan anotasi, stratified random sampling, dan preprocessing berupa penghapusan karakter spesifik, mengubah semua huruf uppercase menjadi lowercase, dan normalisasi lexical serta elongated words. Sehingga diperoleh 9440 data yang digunakan sebagai gold standart dataset. Sehingga data dalam dataset ini adalah data teks *tweet* mengenai RUU Cipta Kerja yang sudah dilabeli “PRO” untuk menunjukkan penulis *tweet* mendukung UU Cipta Kerja, “ANTI” untuk menunjukkan penulis *tweet* menolak UU Cipta Kerja, “ABS” untuk penulis yang netral terhadap UU Cipta Kerja, dan “IRR” untuk *tweet* yang memiliki kata kunci mengenai UU Cipta Kerja namun tidak relevan dengan keberpihakan penulisnya yaitu mendukung, menolak, atau netral.

3. Data Preparation

Data preprocessing yang dilakukan sama dengan teknik-teknik preprocessing yang dilakukan pada [5] yaitu penghapusan karakter spesifik, mengubah huruf kapital menjadi lowercase, normalisasi secara lexical dan kata-kata elongated. Hanya saja diberikan tambahan teknik pengkodean label menggunakan fungsi LabelEncoder dari Sklearn. *Tweet* direpresentasikan menjadi fitur menggunakan pembobotan TF-IDF yang menghasilkan fitur Unigram, bigram, dan fitur kombinasi Unigram+bigram.

Pada model terbaik yang diperoleh yaitu model yang menggunakan klasifier Logistic Regression dengan fitur yang diekstrak menggunakan metode TF-IDF dengan kombinasi Unigram+bigram. Dengan menggunakan classification report yang terlihat pada gambar 1 dapat diketahui bahwa ternyata terdapat penurunan micro f1 score model pada test-set jika dibandingkan dengan training-set. Salah satu penyebabnya adalah penurunan f1-score untuk data berlabel IRR yang sangat signifikan yaitu dari 82.5% menjadi hanya 57%.

Train-set classification report:				
	precision	recall	f1-score	support
ABS	0.842	0.740	0.788	1378
ANTI	0.858	0.966	0.909	3533
IRR	0.920	0.747	0.825	1179
PRO	0.916	0.880	0.898	1462
accuracy			0.874	7552
macro avg	0.884	0.834	0.855	7552
weighted avg	0.876	0.874	0.872	7552

Test-set classification report:				
	precision	recall	f1-score	support
ABS	0.554	0.458	0.502	345
ANTI	0.757	0.888	0.817	884
IRR	0.688	0.486	0.570	294
PRO	0.752	0.738	0.745	366
accuracy			0.718	1889
macro avg	0.688	0.643	0.658	1889
weighted avg	0.708	0.718	0.707	1889

Gambar 1. Classification Report Model Penelitian Terdahulu.

Ditemukan bahwa pada saat pembentukan dataset, Coefficient of Cohan's Kappa Agreement yang diperoleh adalah sebesar 54%. Anotasi dilakukan oleh dua orang annotator yang berbeda, kemudian dataset tersebut dianotasi kembali oleh annotator ketiga dan menggunakan majority vote untuk menyelesaikan perbedaan Annotation Agreement. Ditemukan juga bahwa dari 46% anotasi yang tidak disetujui oleh kedua anotator, dan ternyata terdapat 1473 data atau 15% data berlabel IRR dalam gold standart dataset yang digunakan. Dengan demikian dapat disimpulkan bahwa masing-masing annotator tidak memiliki pemahaman yang baik dan sama terhadap teks *tweet* yang dianggap tidak relevan dengan UU Cipta Kerja atau *tweet* berlabel IRR.

Dengan demikian pada penelitian ini diputuskan untuk tidak menggunakan data yang berlabel IRR, karena annotator tidak dapat melabeli teks yang tidak relevan dengan UU Cipta Kerja. Sehingga dataset yang digunakan berisi 7968 baris data *tweet* berlabel “PRO”, “ANTI”, dan “ABS”. Dengan menggunakan metode train test split, dataset akan dibagi menjadi train-set dan test set dengan perbandingan 80:20. Sehingga 6374 data digunakan untuk train-set dan 1594 digunakan sebagai test-set.

4. Modeling dan Evaluation

Dilakukan inisiasi dari ketiga algoritma klasifikasi yaitu Support Vector Machine, Logistic Regression, dan Multinomial Naïve Bayes. Setelah itu ditetapkan hyperparameter yang akan dilakukan penyesuaian beserta nilai-nilainya untuk masing-masing algoritma.

Pada penelitian ini kami menggunakan metode grid search dengan menggunakan library pipeline dan gridsearchcv dalam membangun model dan menentukan model terbaik. Metode evaluasi yang digunakan adalah K-Fold Cross Validation yang merupakan metode evaluasi default ketika menggunakan gridsearchcv. Gridsearchcv kemudian diinisiasi dengan menggunakan list *classifier*, list parameter, nilai 10 untuk K dalam K-Fold Cross Validation, dan menggunakan f1-score sebagai metrik evaluasi. Menggunakan fungsi gridsearch.fit, dilakukan



pelatihan untuk data trainset. Digunakan fungsi `best_params_` dan `best_scores_` untuk menemukan nilai parameter terbaik dan model terbaik.

```
# SVM Hyperparameters
param1 = {}
param1['classifier'] = [clf1]
param1['extractor__ngram_range'] = [(1,1),(2,2),(1,2)]
param1['extractor__max_features'] = [5000,10000]
param1['classifier__C'] = [10**-2, 10**-1, 10**0, 10**1, 10**2]

# Logistic Regression Hyperparameters
param2 = {}
param2['classifier'] = [clf2]
param2['extractor__ngram_range'] = [(1,1),(2,2),(1,2)]
param2['extractor__max_features'] = [5000,10000]
param2['classifier__multi_class'] = ['ovr','multinomial']
param2['classifier__C'] = [0.001, 0.01, 0.1, 1.0, 10.0]
param2['classifier__solver'] = ['newton-cg', 'lbfgs']
param2['classifier__penalty'] = [None,'l2']

# Multinomial Naive Bayes Hyperparameters
param3 = {}
param3['classifier'] = [clf3]
param3['extractor__ngram_range'] = [(1,1),(2,2),(1,2)]
param3['extractor__max_features'] = [5000,10000]
param3['classifier__alpha'] = [10**0, 10**1, 10**2]
```

Gambar 2. Hyperparameter untuk masing-masing *classifier*.

Dalam upaya meningkatkan performa model, maka akan digunakan parameter-parameter pada gambar 2 dalam eksperimen yang dilakukan. Dapat dilihat bahwa terdapat beberapa perbedaan *Hyperparameter* untuk masing-masing algoritma yang digunakan. Hal ini dikarenakan setiap algoritma memiliki kebutuhan input dan parameter yang berbeda dalam perhitungan yang dilakukan.

III. HASIL DAN PEMBAHASAN

Dari eksperimen yang telah dilakukan diperoleh 288 model. Model-model tersebut dievaluasi menggunakan 10 Fold Cross Validation sehingga dilakukan 2880 fits dalam prosesnya. Setelah itu model-model tersebut diranking sehingga diperoleh lima model terbaik seperti pada tabel 1 berikut:

Tabel 1. Model-model terbaik

Id	Hyperparameter	Mean_test_score	rank
21	'classifier': SVC, C: 10, 'TF-IDF max_features': 10000, 'TF-IDF ngram_range': (1, 1)	0,79181754	1
27	'classifier': SVC, 'C': 100, 'TF-IDF max_features': 10000, 'TF-IDF ngram_range': (1, 1)	0,79166055	2
18	'classifier': SVC, 'C': 10, 'TF-IDF max_features': 5000, 'TF-IDF ngram_range': (1, 1)	0,78993469	3
24	'classifier': SVC, 'C': 100, 'TF-IDF max_features': 5000, 'TF-IDF ngram_range': (1, 1)	0,78977721	4
237	'classifier': LogisticRegression (max_iter=1000), 'C': 10.0, 'multi_class': 'ovr', 'penalty': 'l2', 'solver': 'newton-cg', 'TF-IDF max_features': 10000, 'TF-IDF ngram_range': (1, 1)	0,78883554	5

Perangkingan model-model yang dihasilkan berdasarkan nilai micro f1-score yang dihasilkan. Dalam tabel 1 dapat dilihat terdapat empat model yang menggunakan algoritma SVM sebagai *classifier* yang menempati ranking 1 sampai 4 Terdapat sebuah model yang menduduki urutan kelima dan tidak ada model yang menggunakan algoritma Multinomial Naïve Bayes yang masuk ke dalam ranking lima besar model terbaik.

Dapat diperhatikan juga, bahwa untuk model ranking 1 sampai 4 perbedaan nilai micro f1 score dikarenakan adanya perbedaan nilai pada parameter C pada *classifier* SVM, dan max_features pada teknik ekstraksi fitur menggunakan TF-IDF. Selain itu tidak ada perbedaan nilai f1-score yang signifikan diantara lima model terbaik tersebut.

Model terbaik adalah model_21 yang menggunakan hyperparameter '*classifier*': SVC, C: 10, 'TF-IDF max_features': 10000, dan 'TF-IDF ngram_range': (1, 1). Model tersebut memiliki nilai micro f1-score sebesar 79%. Namun nilai tersebut diperoleh hanya menggunakan train-set dan evaluasi 10 Fold Cross Validation. Untuk itu langkah selanjutnya adalah model terbaik yaitu model_21 tersebut akan diuji menggunakan test-set yang sudah dibagi sebelumnya.

F1_micro Score model_21: 0.7841907151819323				
	precision	recall	f1-score	support
0	0.6406	0.5217	0.5751	345
1	0.8216	0.9026	0.8602	883
2	0.7959	0.7459	0.7701	366
accuracy				
macro avg	0.7527	0.7234	0.7351	1594
weighted avg	0.7765	0.7842	0.7778	1594

Gambar 3. Classification report dari model terbaik

Pada gambar 3 dapat dilihat metrik hasil pengujian model_21 yang diperoleh menggunakan *classification_report*. Pada pemodelan klasifikasi multiclass, metrik micro f1-score nilainya sama dengan precision, recall, dan akurasi dari model. Sehingga dapat dilihat bahwa nilai micro f1-score dari model_21 adalah 78.4%. Dapat diperhatikan bahwa nilai micro f1-score pada train-set hanya turun sebesar 0.7 % jika dibandingkan dengan nilai micro f1-score pada test-set.

IV. KESIMPULAN

Penelitian ini berhasil membangun sebuah model untuk memprediksi keberpihakan penulis teks *tweet* dengan mengklasifikasikannya ke dalam kategori ABS, ANTI, dan PRO. Pada penelitian ini model terbaik memiliki nilai micro f1-score sebesar 78,4%, menggunakan SVM dengan hyperparameter C: 10, 'TF-IDF max_features': 10000, 'TF-IDF ngram_range': (1, 1). Namun pada penelitian ini model yang menggunakan SVM secara umum memiliki performa lebih baik, disusul oleh model yang menggunakan Logistic Regression, dan model yang memiliki performa terendah dihasilkan dengan menggunakan algoritma Multinomial Naïve Bayes. Model terbaik pada penelitian ini termasuk ideal



karena perbedaan micro-f1 score pada train-set dan pada test-set hanya memiliki selisih nilai yang sedikit. Artinya model dapat dikatakan tidak *overfitting* [15].

Metode *hyperparameter tuning* ditambah dengan pengurangan data berlabel IRR pada dataset terbukti dapat meningkatkan performa model. Terbukti adanya peningkatan nilai micro f1-score pada penelitian [5] sebanyak 6,6%.

V. REFERENSI

- [1] M. Yasin, "Mengenal Metode 'Omnibus Law,'" 2020.
<https://www.hukumonline.com/berita/a/mengenal-metode-omnibus-law-lt5f7ad4c048f87/> (accessed May 09, 2023).
- [2] DPRI RI, "Program Legislasi Nasional," 2020.
<https://www.dpr.go.id/uu/prolegnas-long-list> (accessed May 09, 2023).
- [3] M. A. Muqsith, "UU Omnibus law yang Kontroversial," *ADALAH Bul. Huk. KEADILAN*, vol. 4, no. 3, pp. 109–115, 2020, doi: 10.15408/adalah.v4i3.17926.
- [4] Humas MKRI, "MK: Inkonstitusional Bersyarat, UU Cipta Kerja Harus Diperbaiki dalam Jangka Waktu Dua Tahun," 2021.
<https://www.mkri.id/index.php?page=web.Berita&id=17816> (accessed May 09, 2023).
- [5] A. H. Nababan, R. Mahendra, and I. Budi, "Twitter stance detection towards Job Creation Bill," *Procedia Comput. Sci.*, vol. 197, no. 2021, pp. 76–81, 2021, doi: 10.1016/j.procs.2021.12.120.
- [6] D. Küçük and C. A. N. Fazli, "Stance detection: A survey," *ACM Comput. Surv.*, vol. 53, no. 1, 2020, doi: 10.1145/3369026.
- [7] A. ALDayel and W. Magdy, "Stance detection on social media: State of the art and trends," *Inf. Process. Manag.*, vol. 58, no. 4, p. 102597, 2021, doi: 10.1016/j.ipm.2021.102597.
- [8] R. Jannati, R. Mahendra, C. W. Wardhana, and M. Adriani, "Stance Classification Towards Political Figures on Blog Writing," *Proc. 2018 Int. Conf. Asian Lang. Process. IALP 2018*, no. November, pp. 96–101, 2019, doi: 10.1109/IALP.2018.8629144.
- [9] S. Ghosh, P. Singhania, S. Singh, K. Rudra, and S. Ghosh, "Stance Detection in Web and Social Media: A Comparative Study," in *Experimental IR Meets Multilinguality, Multimodality, and Interaction*, vol. 1, Springer, 2019, pp. 75–87. doi: 10.1007/978-3-030-28577-7.
- [10] K. M. Hana, Adiwijaya, S. Al Faraby, and A. Bramantoro, "Multi-label Classification of Indonesian Hate Speech on Twitter Using Support Vector Machines," *2020 Int. Conf. Data Sci. Its Appl. ICoDSA 2020*, 2020, doi: 10.1109/ICoDSA50139.2020.9212992.
- [11] K. S. Chong and N. Shah, "Comparison of Naive Bayes and SVM Classification in Grid-Search Hyperparameter Tuned and Non-Hyperparameter Tuned Healthcare Stock Market Sentiment Analysis," *Int. J. Adv. Comput. Sci. Appl.*, vol. 13, no. 12, pp. 90–94, 2022, doi: 10.14569/IJACSA.2022.0131213.
- [12] A. Shiddicky and S. Agustian, "Analysis of public sentiment against the covid-19 vaccination policy on twitter social media using the logistic regression method," vol. 3, no. 2, pp. 99–106, 2022.
- [13] C. G. Siji George and B. Sumathi, "Grid search tuning of hyperparameters in random forest classifier for customer feedback sentiment prediction," *Int. J. Adv. Comput. Sci. Appl.*, vol. 11, no. 9, pp. 173–178, 2020, doi: 10.14569/IJACSA.2020.0110920.
- [14] M. Swamynathan, *Mastering Machine Learning with Python in Six Steps*, Second Edi. Bangalore: Apress, 2019. doi: <https://doi.org/10.1007/978-1-4842-4947-5>.
- [15] S. Klosterman, *Data Science Projects with Python: A case study approach to gaining valuable insights from real data with machine learning*, Second Edi. Birmingham: Packt Publishing, 2021.