

Sentiment Analysis RKUHP Pada Twitter Menggunakan Metode Support Vector Machine

Tugas Akhir

**diajukan untuk memenuhi salah satu syarat
memperoleh gelar sarjana**

dari Program Studi Informatika

Fakultas Informatika

Universitas Telkom

1301160243

Indera Ihsan



Program Studi Sarjana Informatika

Fakultas Informatika

Universitas Telkom

Bandung

2021

Sentiment Analysis RKUHP Pada Twitter Menggunakan Metode Support Vector Machine

Indera Ihsan¹, Dade Nurjanah², Hani Nurrahmi³

^{1,2,3}Fakultas Informatika, Universitas Telkom, Bandung

⁴Divisi Digital Service PT Telekomunikasi Indonesia

¹inderaihsan@students.telkomuniversity.ac.id, ²dadenurjanah@telkomuniversity.ac.id,

³haninurrahmi@telkomuniversity.ac.id

Abstrak

RKUHP atau rancangan undang-undang kitab hukum dan pidana menuai banyak kritik Indonesia karena dianggap *over criminalization*. Kritik yang disampaikan mayoritas menggunakan media jejaring sosial. Tugas akhir ini menganalisis sentimen terhadap RKUHP dengan melakukan pengklasifikasian terhadap sentimen positif, negatif, dan netral terhadap data yang dikumpulkan dari media sosial Twitter mengenai RKUHP. Penelitian ini dilakukan menggunakan pendapat-pendapat yang disampaikan oleh pengguna jejaring sosial di Indonesia. Metode yang digunakan adalah SVM atau support vector machine dengan mengacu pada penelitian-penelitian sebelumnya yang menunjukkan bahwa SVM memberikan akurasi tertinggi. Data yang digunakan pada penelitian ini adalah tweet pada periode september-november 2019. Pelabelan data dilakukan dengan menggunakan metode *crowd-sourcing* dimana hasil akhir dari label berupa mayoritas hasil pelabelan dari data. Penelitian ini dilakukan dengan memberikan bobot pada setiap data dengan *TF-IDF* dan *sentiment dictionary*, lalu membuat model *machine learning* dengan data yang telah diberikan bobot tersebut. Hasil evaluasi model *machine learning* menggunakan *cross validation* dengan nilai K sebesar 10 serta menggunakan *mean approach* menunjukkan bahwa model memberikan hasil akurasi terbaik sebesar 95% menggunakan kernel *radial basis function*, $C=1000$ dan $\gamma=0.0001$.

Kata kunci : analisis sentimen, support vector machine, prediksi, klasifikasi

Abstract

RKUHP (Rancangan undang-undang kitab hukum pidana) renewal of criminal law was making a huge controversy because considered to have an *over criminalization* value. The critics were mostly given in microblog social media. This research will be done by using the data collected from the social media users in Indonesia about the given topics to retrieve an information about the the sentiment of Indonesian people towards RKUHP. The purpose of this research was to make a classification model which will classify the sentiment from the collected data into three classes : positive, neutral, and negative. This research will also aim to evaluate the performance of the model. The data that will be used in this research is a tweet from the period of september to november 2019. The labelling process in this research was done by crowdsourcing method. In which the majority result of the label will be set as the label. All the labelled data will be weighted using *TF-IDF* and *sentiment dictionary* and later on the weighted matrix will be used to build a machine learning model. The evaluation result of the machine using *cross validation* with the value of K equals to 10 using *mean approach* shows that the model reaches the highest accuracy with 95% using *radial basis function* kernel, $C=1000$ and $\gamma=0.0001$.

Keywords: sentiment analysis, support vector machine, prediction, classification

1. Pendahuluan

1.1 Latar Belakang

Dewasa ini, penggunaan media sosial sedang berkembang pesat. Perusahaan, organisasi, masyarakat menggunakan media sosial untuk mendapatkan informasi mengenai pandangan dan perasaan masyarakat terhadap perusahaan atau organisasi tersebut[1]. Salah satu media sosial yang kerap digunakan adalah Twitter[2]. Twitter adalah sebuah layanan microblogging dan jejaring sosial yang memungkinkan penggunanya untuk memposting pesan yang disebut sebagai tweet dan berinteraksi dengan pengguna lainnya [3]. Di Indonesia, Twitter memiliki 19,5 juta pengguna dari total 500 juta pengguna global. Hal ini memungkinkan terjadinya perbedaan ungkapan dari masyarakat dengan latar belakang yang berbeda terhadap sebuah persoalan yang terjadi[2][3]. Informasi mengenai pro dan kontra tersebut dapat digunakan sebagai deskripsi pandangan masyarakat terhadap sebuah masalah[4]. Rancangan Undang-Undang Kitab Undang-Undang Dan Hukum Pidana (RUU KUHP) yang saat ini belum disahkan menuai banyak kritik di Indonesia karena dianggap *over criminalization* [5]. Banyaknya kritik yang disampaikan di Twitter menyebabkan #RKUHP menjadi salah satu topik yang menjadi Trending atau topik yang paling sering dibahas tahun 2019. identifikasi pro dan kontra terhadap RKUHP dengan melakukan analisis

dan klasifikasi terhadap cuitan atau tweet terhadap RKUHP dapat menghasilkan sebuah representasi dari opini masyarakat terhadap kasus tersebut. salah satu metode yang digunakan sebelumnya dan memberikan akurasi terbaik untuk analisis sentimen adalah Support Vector Machine(SVM) [6]. Sampai saat ini belum ada sebuah tools untuk mengukur sentimen masyarakat Indonesia terhadap RKUHP. Maka dari itu hasil analisis akan memberikan sebuah informasi representatif terhadap perasaan dan pendapat pengguna twitter di Indonesia mengenai RKUHP.

1.2 Topik dan Batasannya

Sampai saat ini, belum ada sebuah informasi yang bersifat representatif akan pendapat dan perasaan rakyat Indonesia mengenai revisi kitab hukum undang undang dan pidana (RKUHP) yang dianggap bersifat over criminalization dan represif terhadap rakyat Indonesia [4] . Sebuah informasi yang bersifat representatif dapat digunakan untuk menjabarkan sebuah alasan mengapa RKUHP menuai banyak kritik di media sosial. Maka dari itu didapatkan beberapa kesimpulan mengenai permasalahan yang terdapat saat ini adalah :

A. Belum adanya sebuah informasi yang bersifat representatif terhadap opini masyarakat mengenai RKUHP

B. Belum didapatkan sebuah tools untuk mengukur sentimen masyarakat Indonesia mengenai RKUHP

Berdasarkan pemaparan tersebut maka didapatkan rumusan masalah pada penelitian ini : Bagaimana mengukur sentiment analysis berdasarkan perasaan dan opini masyarakat Indonesia mengenai RKUHP?

Pada penelitian-penelitian sebelumnya [14][15] memberikan hasil bahwa metode sentiment analisis yang memberikan matriks akurasi paling baik adalah Support Vector Machine dan Multinomial Naïve Bayes. Namun setelah dilakukan proses preprocessing dan menggunakan dataset mengenai topik politik seperti pada penelitian ini, didapatkan sebuah informasi bahwa Support Vector Machine memberikan hasil akurasi yang lebih baik. Data yang digunakan merupakan hasil penambangan data dari twitter yang memiliki dua komponen, yaitu target dan sentimen, dimana target merupakan sebuah entitas yang menjadi tujuan dari opini dan sentiment (opini) adalah sebuah pernyataan atau ekspresi yang bernilai positif, negative, atau netral seperti pada penelitian sebelumnya [7][4]. Data yang didapatkan sebagai data latih adalah tweets antara bulan September sampai Desember tahun 2019. Penentuan sentimen dilakukan dengan metode crowd sourcing. Dengan meninjau kata yang mengandung opini baik yang memiliki polarity positif maupun negatif dari tweet yang sudah dilabeli kelas katanya. Kelas kata yang dipilih adalah kata sifat (adjective), kata keterangan (adverb), kata benda (noun) dan kata kerja (verb), sesuai dengan penelitian [4]. Data latih yang diambil pada twitter adalah data latih berupa tweets dengan tagar #RKUHP.

1.3 Tujuan

Tujuan penelitian yang dilakukan adalah untuk mengukur sentimen masyarakat Indonesia terhadap RKUHP berdasarkan data yang didapat pada media jejaring social *Twitter*

Organisasi Tulisan

Sistematika penulisan pada penelitian ini dilakukan dengan bab pertama membahas mengenai permasalahan yang pada penelitian ini yaitu meliputi latar belakang, rumusan permasalahan, dan tujuan. Selanjutnya pada bab kedua merupakan segala informasi yang dijadikan acuan dalam menyelesaikan permasalahan pada bab pertama baik berupa penelitian sebelumnya maupun paper lain yang terkait. Pada bab ketiga merupakan perancangan sistem yang dilakukan, hasil dari ekstraksi informasi pada bab kedua. Setelah bab ketiga, hasil dari sistem yang dibuat akan evaluasi dengan beberapa matriks pada bab keempat. Terakhir pada bab kelima adalah kesimpulan dan saran pada penelitian yang sudah dilakukan.

2. Studi Terkait

2.1 Analisis Sentimen

Analisis sentiment adalah sebuah implementasi dari pemrosesan bahasa alami yang mempelajari opini, sentiment, dan emosi yang di ekspresikan dalam sebuah teks[8]. Untuk memperjelas definisi analisis sentiment diberikan sebuah contoh review dari sebuah buku:

“Saya kemarin membeli buku Albert Camus (1) dalam buku tersebut didapatkan sebuah informasi yang menarik (2) meskipun cover dari buku tersebut kurang menarik (3) kualitas kertasnya bagus (4) itu saya menyukai buku itu. (5) “

Terdapat beberapa opini pada kalimat diatas. Kalimat (2), (4),(5) mengekspresikan sebuah opini positif. Sedangkan kalimat (3) mengekspresikan sebuah opini negative. Kemudian ada yang disebut sebagai target, target yang dimaksud adalah sebuah tujuan atau objek dari opini atau emosi tersebut. pada kalimat (2) objek yang menjadi target dari opini adalah informasi, pada kalimat (3) target yang menjadi objek dari opini adalah cover buku, pada kalimat (4) targetnya adalah kualitas kertas.secara umum opini dapat diekspresikan mengenai apapun seperti

produk, organisasi, aktivitas sosial, olahraga dan lain-lain. Berdasarkan definisi dari analisis sentiment, maka dapat disimpulkan bahwa opini memiliki dua komponen yaitu target g dan sentiment s , dimana g adalah aspek dari opini tersebut dan bisa berupa apapun sementara s bernilai positif, negative atau netral [8]. Sebagai contoh pada sub-bab sebelumnya target dari opini pada kalimat adalah “Buku Albert Camus” dan sentimen yang diberikan adalah Informasi pada kalimat (1). Serta terdapat atribut dari topik seperti pada kalimat (3) terdapat atribut dari topik yaitu “cover buku” yang dapat didefinisikan sebagai “o” pada review tersebut terdapat subjek yang memaparkan opininya yaitu “aku” dan waktu review tersebut dikeluarkan adalah 1 November 2019. Dapat didefinisikan bahwa sebuah opini adalah quadruple (g,o,s,h,t) g adalah target atau topik dari opini, o merupakan atribut dari topic, s merupakan sentiment yang bisa bernilai positif, negative atau netral, h merupakan subjek yang memberikan opini dan t adalah waktu saat opini diekspresikan. Berdasarkan hasil penelitian sebelumnya dan studi literature yang didapatkan sebuah tahapan untuk melakukan Analisis Sentimen[4] [8] sebuah kesatuan yang didefinisikan dengan D sebagai contoh pada kalimat (3) didapatkan sebuah kesatuan opini seperti contoh :

$$D(3) = ("Buku Albert Camus" _ "Cover Buku" _ Negative _ "Saya" _ "1-10-2019")$$

Pada penelitian ini analisis sentiment dilakukan dengan klasifikasi.

2.2 Corpus

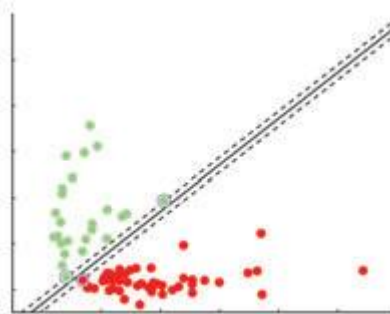
Corpus merupakan sekumpulan kata dalam satu bahasa yang terbentuk menjadi satu kesatuan. Pada penelitian ini korpus yang digunakan adalah *Inset Lexicon* yang dibentuk pada penelitian terkait [14][15] dimana corpus pada penelitian ini akan digunakan untuk memberikan bobot pada setiap kata yang terdapat pada kalimat.

	word	score
0	hai	3
1	merekam	2
2	ekstensif	3
3	paripurna	1
4	detail	2

Tabel 1. Corpus dan Score

2.3 Support Vector Machine

Support Vector Machine adalah sebuah algoritma yang melabeli sebuah objek dengan mempelajari sampel data contoh [9]. Yang membuat support vector machine berbeda adalah pencarian hyperplane terbaik sebagai pemisah yang tidak terdapat pada mesin pembelajaran lainnya[10]. Konsep kerja Support Vector Machine adalah upaya untuk mencari hyperplane terbaik untuk memisahkan dua buah kelas pada input space dengan memberikan label +1 atau -1



Gambar 1. Hyperplane pada SVM

Gambar menunjukkan terklasifikasinya dua buah kelas dengan label +1 berwarna hijau dan -1 berwarna merah [9]. Usaha untuk mencari lokasi hyperplane ini merupakan tujuan dari proses pembelajaran pada SVM. Dalam SVM terdapat beberapa jenis kernel yaitu kernel gaussian, kernel linear, dan kernel polynomial. Pada

Gambar memperlihatkan beberapa pattern yang merupakan anggota dari dua buah kelas, yaitu kelas +1 dan kelas -1. Pattern yang tergabung pada kelas -1 disimbolkan dengan warna merah (kotak), sedangkan pattern pada kelas +1 disimbolkan dengan warna kuning (lingkaran). Problem klasifikasi dapat diterjemahkan dengan usaha untuk menemukan garis hyperplane terbaik yang memisahkan antara kedua kelas tersebut[10]. Hyperplane pemisah terbaik di antara kedua kelas dapat ditemukan dengan mengukur margin dari hyperplane tersebut serta mencari titik maksimalnya. Dalam SVM terdapat beberapa jenis kernel yaitu kernel *gaussian*, kernel *polynomial* dan kernel *Radial Basis Function*. Dengan perhitungan formula seperti yang disajikan dalam Tabel sebagai berikut:

Tabel 2 Kernel Of SVM

Jenis	Fungsi
Radial Basis Function	$K(x_i, x_j) = \exp\left(-\frac{\ x_i - x_j\ ^2}{2\sigma^2}\right)$
Polynomial	$K(x_i, x_j) = ((x_i \cdot x_j) + c)^d$
Linear	$K(x_i, x_j) = (x_i \cdot x_j) + c$

2.4 Implementasi Support Vector Machine Classifier Untuk Menentukan Polaritas Sentimen

Setelah mendapatkan hasil berupa nilai sentimen pada data yang didapatkan, maka tahap selanjutnya adalah menciptakan tools untuk mengklasifikasi polaritas dari sentiment sebuah kalimat dengan memberikan label pada setiap kalimat yang akan diprediksi. Berikut merupakan sebuah penyesuaian dari support vector machine berdasarkan penelitian [8] : Biarkan $X \rightarrow Y$ menjadi pemetaan. Di sini X berkorespondensi untuk sentimen yang diekstraksi dari sebuah kalimat. X memiliki nilai label berupa skor positif atau negatif. Skor ini dapat ditentukan secara manual seperti pada pendefinisian analisis sentiment. Kasus klasifikasi paling sederhana dan digunakan untuk pengklasifan sentiment adalah klasifikasi biner. Namun pada penelitian ini klasifikasi menambahkan kelas ketiga yaitu kelas polaritas netral. pemetaan $X \rightarrow Y$ adalah sebuah nilai pasangan yang dimana $Y \in R$. Y dapat mengambil polaritas negatif dan positif. Pelatihan sistem dapat dilakukan menggunakan kumpulan fitur $x_1: y_1, x_2: y_2, \dots, x_m: y_m$. Setelah proses pelatihan, mesin pembelajaran dapat mempelajari classifier dan pemisahan pendapat positif dan negatif akan dapat dilakukan oleh SVM. Penelitian ini dilakukan menggunakan mesin SVM dengan tipe kernel radial basis function. Semua parameter dari classifier diatur ke nilai standarnya.

2.5 Optimisasi Hyperparameter menggunakan Grid Search

Secara umum algoritma machine learning tidak akan memberikan performa yang baik apabila tidak dilakukan pengaturan parameter terhadap algoritma machine learning tersebut maka dari itu sangat penting untuk melakukan optimisasi parameter terhadap algoritma machine learning yang akan dibangun [18]. Optimisasi parameter akan membutuhkan waktu yang sangat lama terutama jika algoritma machine learning yang dibangun memiliki banyak parameter[18][12] pada penelitian ini metode optimisasi yang digunakan adalah *grid search*. *grid search* adalah sebuah pencarian lengkap dengan mengombinasikan seluruh hyperparameter yang diberikan. Hyperparameter didefinisikan dengan nilai minimum, nilai maximum dan skala antar nilai didalamnya

Grid search mengoptimasi parameter dari SVM (C , *kernel*, γ , dsb.) menggunakan *cross validation technique* sebagai matriks performansi. Tujuannya adalah untuk menemukan kombinasi terbaik dari *hyper-parameter* sehingga model klasifikasi dapat memprediksi data yang belum diketahui secara akurat. [18]. terdapat dua parameter yang sangat berpengaruh pada performa dari SVM yaitu parameter regularisasi C , yang menentukan tradeoff antara meminimalkan kesalahan pelatihan dan meminimalkan kompleksitas model dan parameter *gamma* dari fungsi kernel yang secara implisit memisahkan pemetaan nonlinier dari ruang masukan ke ruang fitur berdimensi tinggi.

2.6 Skenario Evaluasi Model

Untuk mengevaluasi performansi dari klasifikasi yang sudah dilakukan, maka diperlukan perhitungan accuracy, precision, recall, dan F1 score. Confusion Matrix sering digunakan dalam machine learning. Confusion Matrix digunakan untuk mengevaluasi performansi dari metode klasifikasi dengan dataset yang digunakan. Struktur dari confusion matrix di representasikan kedalam baris dan kolom, dimana baris adalah data aktual dan kolom adalah data yang diprediksi[11]. mengukur performa dari model SVM. Pada proses evaluasi akan meliputi beberapa langkah yaitu nilai prediksi akurasi, presisi, recall, dan F1 score dari data

Tabel 3. Contoh *Confusion Matrix*

		Data Prediksi		
Data Aktual	Negatif	39	0	0
	Netral	2	8	0
	Positif	0	0	18
		Negatif	Netral	Positif

Pada contoh confusion matrix tersebut terdapat tiga kelas yaitu Negatif, Positif dan Netral. Dimana pada proses tersebut terdapat Label sebenarnya (biasa disebut dengan y_true) dan hasil dari prediksi data test (y_pred) contoh dari y_true dan y_pred untuk Label Netral adalah sebagai berikut

$$Conf(y_true, y_pred) = [2, 8, 0]$$

Dimana 2 adalah false_negatif, 8 adalah true_netral, dan 0 adalah false_positif. Dapat dikatakan dengan kalimat “dari 10 data test berlabel netral (2+8) mesin dapat memprediksi 8 true netral dan 2 false negatif”. Penghitungan akurasi dilakukan dengan cara :

$$Accuracy = \frac{TP+TN+TNet}{N} \times 100$$

Dimana

N=jumlah seluruh data test

TP=true positive

TN=true Negative

Tnet=true netral

Sehingga berdasarkan contoh *confusion matrix* didapatkan akurasi :

$$(39+8+18) \times 100 / 67 = 97.014$$

Evaluasi selanjutnya adalah precision untuk melakukan *precision* dilakukan dengan cara :

$$Precision = \frac{TP}{TP + FP} \times 100$$

Sehingga didapatkan hasil penghitungan *precision* sebagai berikut :

	Positif	Netral	Negative	Total
TP/(TP+FP)	18/18	8/10	39/39	2,08
Result	1	0.8	1	(1+1+0.8)/3=0.93

Evaluasi selanjutnya adalah recall untuk melakukan recall dilakukan dengan cara :

$$Recall = \frac{TP}{TP + FN + FNET} \times 100$$

Sehingga didapatkan nilai Recall untuk setiap kelas sebagai berikut :

	Positif	Netral	Negative	Total
TP/(TP+FN+FNET)	18/18	8/8	39/(39+2)	2,095
Result	1	1	0.95	(1+1+0.95)/3=0.983

Evaluasi selanjutnya dilakukan dengan menghitung nilai F1 dengan cara perhitungan sebagai berikut :

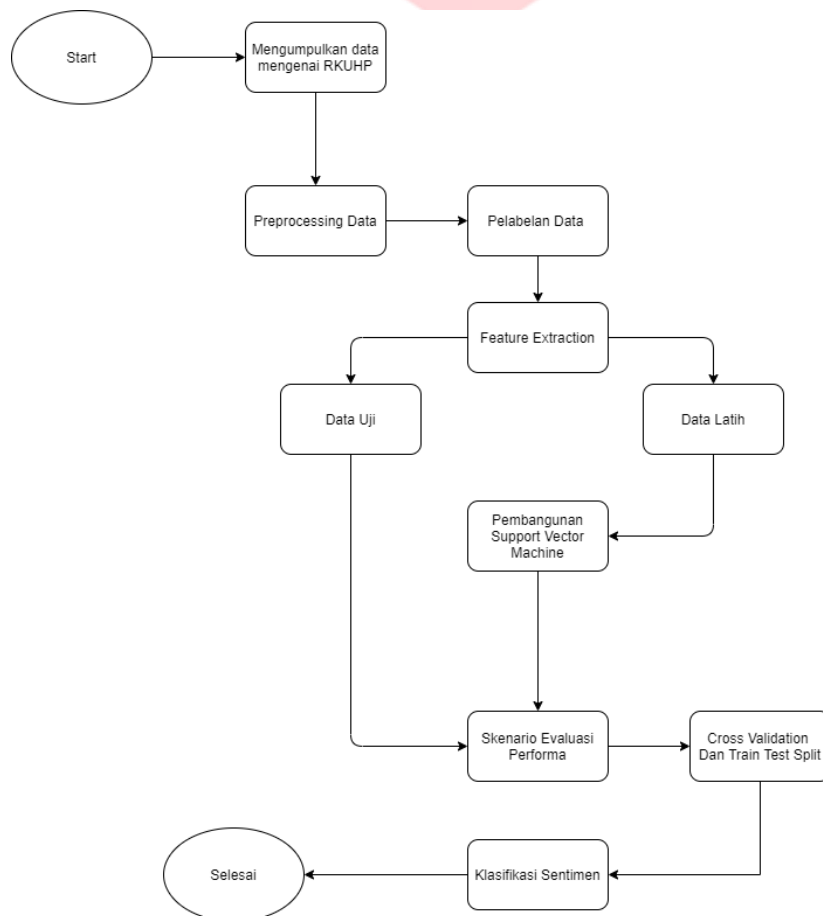
$$F1 = 2x \frac{precision(a) \times recall(a)}{precision(a) + recall(a)}$$

Sehingga didapatkan nilai F1 untuk seluruh kelas sebagai berikut :

	Positif	Netral	Negative	Total
$2x \frac{p(a) \times r(a)}{p(a) + r(a)}$	$2x \frac{1 \times 1}{1+1}$	$2x \frac{0.8 \times 1}{0.8+1}$	$2x \frac{1 \times 0.95}{1+0.95}$	2,86
Result	1	0.89	0.97	$(1+0.89+0.97)/3=0.953$

3. Sistem yang Dibangun

Pada penelitian ini terdapat beberapa tahapan pada sistem sebelum memberikan hasil penelitian. Berikut ini diagram alir yang menjelaskan tahapan perancangan sistem dalam mendapatkan hasil akhir dari penelitian:



Gambar 2. Diagram Alir Perancangan Sistem

3.1 Mengumpulkan data mengenai RKUHP

Pada proses ini dilakukan proses untuk mengumpulkan data yang akan digunakan pada penelitian, data yang dikumpulkan didapat dengan melakukan proses *scrapping* pada twitter dengan filter tagar RKUHP pada rentang bulan September dan bulan November 2019 didapatkan data dalam bentuk *Coma Separated Value* (CSV) sebanyak 2612 data.

Tabel 4. Contoh Sampel Data

No.	Data
1	Komisi III @DPR_RI @bambangsoesatyo, Tolak RKUHP yg kriminalisasi perempuan, anak, masyarakat adat & kelompok marjinal. - Tandatangani Petisi! http://chnng.it/bP6hw289A lewat @ChangeOrg_ID""
2	Pemerintah dan DPR cukup membuat kodifikasi di dalam RKUHP. #TolakRUUKKS #SiberIndonesia #NKRIHargaMatipic.twitter.com/rXrL6uQqSt

3.2 Preprocessing

Pada proses ini dilakukan persiapan data text seperti pengubahan setiap karakter dalam text menjadi huruf Non-kapital, penghapusan *emoticon*, penghapusan tanda baca dan *link* serta membuat teks menjadi dalam bentuk token. Fokus pada tahap ini adalah untuk menghapus noise yang terdapat pada text yang akan digunakan. Tahapan yang dilakukan pada preprocessing yang pertama adalah mengubah seluruh kalimat menjadi dalam bentuk Karakter Non kapital, penghapusan tanda baca, *emoticon* dan URL dengan contoh sebagai berikut

Tabel 5. Perbandingan data preprocessing

No.	Setelah Preprocessing	Sebelum Preprocessing
1.	komisi iii dpr_ri bambangsoesatyo tolak rkuhp yg kriminalisasi perempuan anak masyarakat adat kelompok marjinal tandatangani petisi lewat changeorg_id	Komisi III @DPR_RI @bambangsoesatyo, Tolak RKUHP yg kriminalisasi perempuan, anak, masyarakat adat & kelompok marjinal. - Tandatangani Petisi! http://chnng.it/bP6hw289A lewat @ChangeOrg_ID""
2.	pemerintah dan dpr cukup membuat kodifikasi di rkuhp aja lah indonesiamaju	Pemerintah dan DPR cukup membuat kodifikasi di dalam RKUHP. #TolakRUUKKS #SiberIndonesia #NKRIHargaMatipic.twitter.com/rXrL6uQqSt

3.3 Pelabelan Data

Proses pelabelan data dilakukan dengan tujuan untuk mendapatkan kelas yang akan dijadikan sebagai Label untuk dipelajari oleh algoritme Support Vector Machine. Kelas yang dibentuk terdiri dari tiga kelas yaitu kelas positif negative dan netral. Metode pelabelan dilakukan dengan cara crowdsourcing dengan responden berjumlah 5 orang yang kemudian mayoritas label yang diberikan akan diambil seperti pada contoh berikut :

Responden 1	Responden 2	Responden 3	Responden 4	Responden 5	Hasil Label
Positive	Positive	Positive	Netral	Negatif	Positive

3.4 Feature Extraction

Proses ini merupakan proses untuk mengubah text yang dibentuk kedalam bentuk matrix dengan isi berupa bobot sehingga dapat dipelajari oleh mesin. Fitur yang digunakan pada penelitian kali ini adalah TF, TF-IDF, dan Sentiment Dictionaries.

Fitur pertama yang digunakan pada penelitian ini adalah TF-IDF yang dikenal sebagai suatu ekstraksi fitur yang digunakan untuk memberikan suatu bobot, dan performansinya bahkan masih bisa dibandingkan dengan teknik teknik baru. Dokumen akan digunakan sebagai faktor dalam pembobotan[16]. Pada metode TF-IDF terdapat dua proses, yaitu penghitungan nilai TF (Term Frequency) dan IDF (Inverse Document Frequency). Pada setiap term, data akan diberikan suatu bobot[16]. Untuk menghitung bobot term (t) yang ada pada dokumen (d) dibutuhkan persamaan sebagai berikut :

$$idf_t = \log \frac{N}{df_t}$$

$$w_{t,d} = tf_{t,d} \times idf_t$$

tf = jumlah term dalam dokumen

idf = jumlah term pada seluruh dokumen yang ada pada data dengan persamaan sebagai berikut :

dimana :

df : jumlah dokumen yang didalamnya terdapat term t, N = jumlah seluruh dokumen yang digunakan dalam data.

Fitur kedua yang akan digunakan pada penelitian ini adalah Sentiment Lexicon . Pada penelitian ini sentiment lexicon yang digunakan merupakan hasil penelitian dari [14][15] yang berisi sekumpulan kata positive dan negative yang kemudian akan digunakan sebagai acuan untuk mendapatkan bobot opini dari setiap tweet pembobotan dilakukan dengan mengacu pada setiap sentiment lexicon yang digunakan dengan tujuan untuk mendapatkan bobot nilai dengan persamaan :

$$S = \sum_{i=1}^n P(i) - N(i)$$

Dimana :

S : Nilai dari bobot sebuah bait .

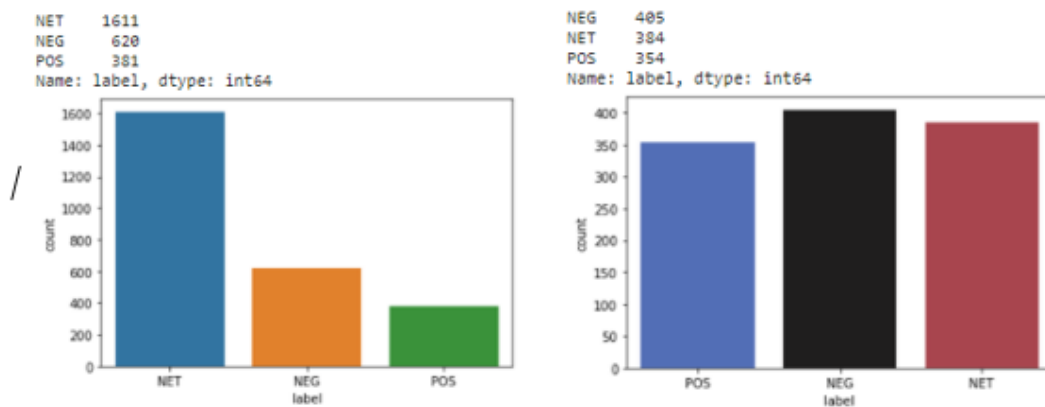
P(i):Bobot positif kata (i) yang terdapat pada korpus

N(i): Bobot Negatif kata(i) Pada korpus

3.5 Data uji dan data latih

Pada tahap ini dilakukan pembagian untuk melakukan pengujian terhadap algoritma SVM. Dimana data yang sudah dilakukan preprocessing dan mendapatkan feature extraction akan dibagi menjadi data latih dan data uji. Pembagian data latih dan data uji dilakukan dengan rasio 80:20 dimana data latih sebanyak 80% dan data uji sebanyak 20% . data latih akan digunakan sebagai data yang akan melatih algoritma SVM dan data uji adalah data yang akan dicoba untuk diprediksi oleh algoritma. Tahap ini juga membandingkan beberapa metode klasifikasi seperti penelitian sebelumnya [14][15][17] untuk mendapatkan informasi mengenai performa tiap metode klasifikasi dengan dataset yang digunakan. Metode klasifikasi yang digunakan pada tahap ini adalah K Nearest Neighbour, Multinomial Naïve Bayes, Support Vector Machine dimana seluruh parameter tiap metode klasifikasi diatur dengan nilai standar.

didapatkan bahwa data yang didapat bersifat imbalanced dimana jumlah setiap kelas memiliki selisih yang besar seperti pada gambar dibawah :

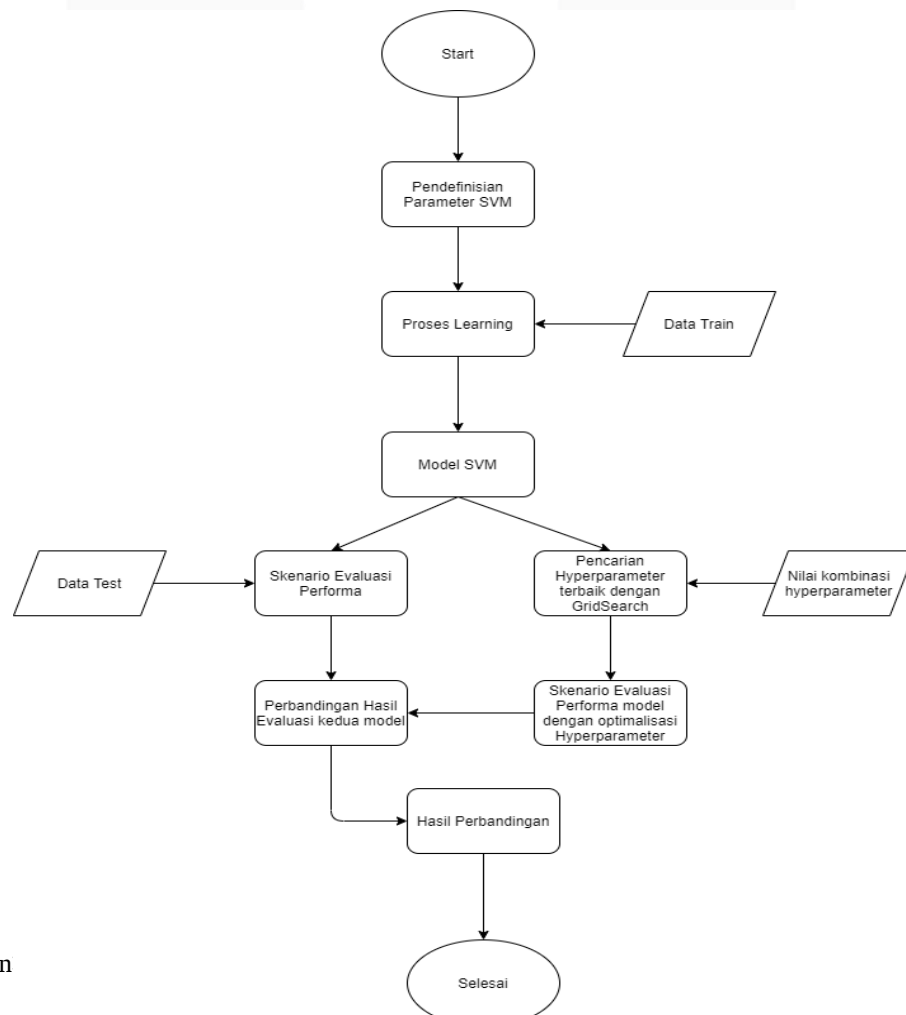


Gambar 3. Perbandingan data sebelum dan sesudah downsampling

Maka dari itu diperlukan proses downsampling untuk kelas NET dan NEG agar mendapatkan data dengan jumlah sample yang seimbang. Downsampling dilakukan dengan cara memilih sample dari kelas NET dan NEG sejumlah sama dengan kelas POS sehingga mendapatkan jumlah yang sama pada setiap kelas. Kemudian pelabelan data dilakukan kembali.

3.6 Pembangunan Support Vector Machine

Setelah data diubah menjadi dalam bentuk matriks pada tahap feature extraction, matriks dan label yang didapatkan pada tahap selanjutnya akan menjadi bahan pembelajaran bagi SVM sehingga SVM dapat memprediksi label dari sebuah matrix yang merupakan fitur yang sama dengan matrix yang telah dipelajari sebelumnya. Cara kerja dari SVM ialah menciptakan sebuah hyperplane yang akan memisahkan data menjadi beberapa kelas.



Gambar 4. Algoritma Pembangunan Model Support Vector Machine

Pada tahap ini parameter SVM yang akan mencoba untuk mempelajari feature diisi dengan nilai standard. kemudian di tahap selanjutnya dilakukan proses learning oleh SVM menggunakan feature dari data test dan label dari data test yang pada tahap sebelumnya telah didapatkan. Setelah Model terbentuk dilakukan evaluasi dengan menggunakan *cross validation* dengan nilai K sebesar 10 untuk mengukur performa dari model yang dibuat dan menjadikan performa tersebut sebagai *baseline* dari penelitian ini. Kemudian dilakukan optimalisasi terhadap parameter dari SVM dengan menggunakan *gridsearch* dan mendapatkan parameter terbaik yang dapat digunakan oleh model SVM. Berikut parameter serta nilai yang digunakan pada metode SVM seperti yang terlihat pada Tabel 6 sebagai berikut:

Tabel 6. Kombinasi Hyperparameter SVM

Parameter	Nilai Parameter
C	1,10,100,1000,10000
Gamma	0.1 , 0.01, 0.0001,0.00001
Kernel	RBF

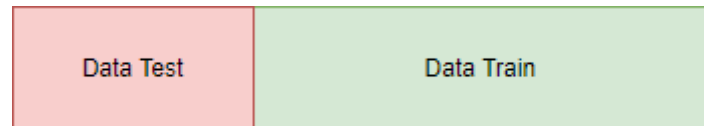
Gridsearch akan mencoba untuk menemukan *hyperparameter* terbaik dengan cara mengevaluasi tiap kombinasi *hyperparameter* dengan menggunakan *K-Fold*. Dan memilah kombinasi terbaik dari nilai *hyperparameter* yang diberikan. Berikut merupakan parameter terbaik dari *gridsearch* :

Tabel 7. Kombinasi Hyperparameter Terbaik Gridsearch

Parameter	Nilai Parameter Terbaik
C	1000
Gamma	0.0001
Kernel	Radial Basis Function

Evaluasi.**4.1 Hasil Pengujian**

Pada tahap pengujian dimana seluruh parameter dari SVM diatur ke nilai standardnya didapatkan bahwa performa dari SVM diukur dengan skenario evaluasi performansi pemisahan data test dan data train dengan rasio 20:80 dengan ilustrasi dan hasil sebagai berikut :

Gambar 5. Ilustrasi *Train Test Split*

Tabel 8. Hasil Skenario Evaluasi SVM dengan Train Test Split

SVM	Kelas	Precision	Recall	F1 Score	Accuracy
	POS	0.80	0.66	0.72	
	NEG	0.74	0.68	0.71	
	NET	0.53	0.68	0.60	

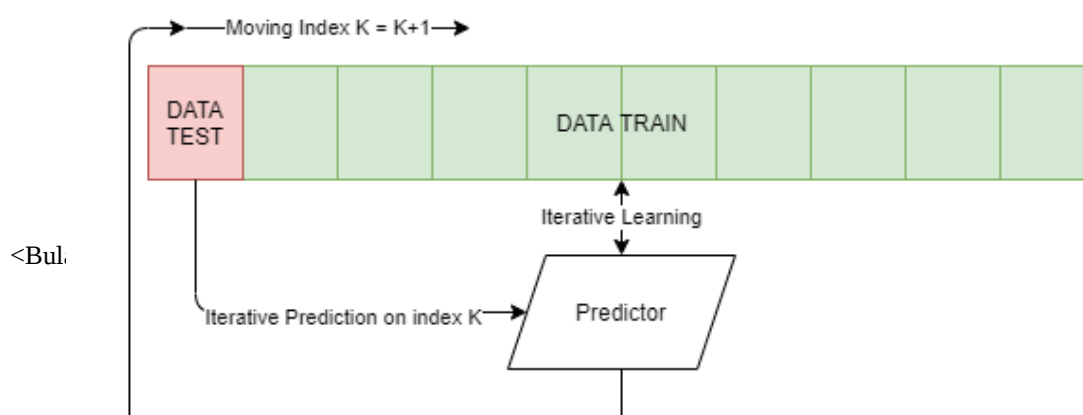
Hasil menunjukkan bahwa algoritma machine learning SVM memiliki performa yang kurang baik . dengan Akurasi sebesar 67% sehingga dapat disimpulkan bahwa optimisasi *hyperparameter* dari SVM perlu dilakukan untuk meningkatkan performa dari SVM . Namun precision pada kelas positive dapat dikatakan cukup baik dengan nilai 0.80 namun pada kelas Netral justru SVM belum memberikan hasil yang baik. Hal ini dapat terjadi karena penggunaan korpus yang hanya mengindikasikan kalimat negative dan positive sehingga machine learning dengan tiga kelas seperti pada penelitian ini hanya dapat berfokus untuk mengklasifikasikan kalimat yang terindikasi sebagai kalimat negative dan kalimat positive. Menurut penelitian sebelumnya [14][15][16] dinyatakan bahwa dalam melakukan klasifikasi untuk sentiment analysis , Multinomial Naïve Bayes dan SVM memberikan hasil yang baik, sehingga pada hasil pengujian hasil performa dari SVM akan disandingkan dengan Multinomial Naïve Bayes

Tabel 9. Hasil Skenario Evaluasi Multinomial Naïve Bayes dengan Train Test Split

Multi Nomial Naïve Bayes	Kelas	Precision	Recall	F1 Score	Accuracy
	POS	0.82	0.87	0.85	
	NEG	0.60	0.97	0.74	
	NET	1.00	0.30	0.46	

Dapat dikatakan bahwa algoritma Multinomial Naïve Bayes memberikan performa yang lebih baik dibandingkan SVM . dengan akurasi sebesar 72% dan nilai precision yang sangat baik sebesar 100% untuk kelas netral . dengan ini dapat dikatakan bahwa dengan mengimplementasikan SVM dan Multinomial Naïve Bayes dengan parameter standard SVM tidak dapat mengungguli multinomial Naïve Bayes dengan skenario evaluasi *train test split*. Sehingga dapat disimpulkan tanpa adanya optimalisasi *hyperparameter* pada SVM , SVM belum dapat mengukur sentiment dari rakyat Indonesia yang tergolong dalam kelas netral terkait dengan RKUHP . namun untuk kalimat yang tergolong dalam kelas positive dan negative SVM memberikan performa yang cukup baik.

Skenario evaluasi kedua adalah dengan menggunakan cross validation yang mana pada evaluasi ini dataset akan dicacah menjadi 10 bagian dan pada tiap bagian data akan digunakan sebagai data test, dan yang lain akan digunakan sebagai data train. Berikut merupakan ilustrasi dari cross validation :



Gambar 6 . Ilustrasi *cross validation*

Dengan *cross validation*, model akan melakukan proses evaluasi secara iteratif dimana index K akan bergerak hingga index terakhir, dan data yang terletak pada partisi K akan menjadi *data test* dan yang lain akan menjadi data train. Model prediksi akan mencoba untuk memprediksi data yang terdapat di index K. Setelah proses prediksi selesai maka hasil akan disimpan menjadi satu kesatuan. Sehingga akan terdapat hasil prediksi sejumlah K. Setelah seluruh hasil prediksi didapatkan dilakukan *mean approach* untuk mendapatkan hasil akhir dengan menghitung rata-rata dari seluruh nilai yang didapat pada setiap iterasi.

Sehingga didapatkan hasil sebagai berikut :

Tabel 10. Hasil Skenario Evaluasi SVM dan *Multinomial Naïve Bayes* dengan *cross validation*

	Index K	Accuracy	Precision	Recall	F1 Score
SVM	1	0.54	0.56	0.55	0.54
	2	0.69	0.73	0.69	0.70
	3	0.66	0.68	0.65	0.66
	4	0.70	0.70	0.69	0.70
	5	0.71	0.74	0.71	0.71
	6	0.71	0.71	0.70	0.70
	7	0.70	0.70	0.70	0.70
	8	0.80	0.80	0.80	0.80
	9	0.75	0.78	0.75	0.75
	10	0.77	0.8	0.76	0.77
rata-rata		0.70	0.72	0.70	0.70

	Index K	Accuracy	Precision	Recall	F1 Score
Multi Nomial Naïve Bayes	1	0.53	0.67	0.54	0.45
	2	0.72	0.80	0.72	0.65
	3	0.68	0.77	0.68	0.63
	4	0.71	0.80	0.71	0.64
	5	0.74	0.81	0.74	0.69
	6	0.69	0.80	0.68	0.63
	7	0.68	0.78	0.68	0.60
	8	0.67	0.77	0.67	0.56
	9	0.66	0.77	0.66	0.58
	10	0.67	0.77	0.68	0.61
rata-rata		0.67	0.77	0.68	0.60

Pada evaluasi *cross validation* dan menggunakan *mean approach* didapatkan bahwa performa dari SVM ternyata memiliki nilai yang lebih baik dibanding dengan Multinomial Naïve bayes dari segi akurasi, recall dan F-1 Score. Namun dari matrix precision, Multinomial Naïve bayes lebih unggul dibanding dengan SVM. Nilai maximum akurasi, precision recall dan F1 Score secara berurutan dari SVM adalah 0.8, 0.8, 0.8 dan 0.8 sedangkan dengan naïve bayes secara berurutan adalah 0.74, 0.81, 0.74, 0.69. Jika menggunakan Optimistic approach, dimana nilai maximum dari setiap matrix pada index K yang diambil maka SVM dapat dikatakan dapat

mengungguli Multinomial Naïve bayes. Namun dengan ini tetap perlu dilakukan *hyperparameter* pada SVM untuk melihat tingkat performa dari SVM ketika sudah dilakukan *hyperparameter optimization*

Tabel 11. Hasil Skenario Evaluasi SVM dengan *parameter tuning* dan *Multinomial Naïve Bayes*

SVM	Kelas	Precision	Recall	F1 Score	Accuracy
	POS	0.94	1.00	0.97	
	NEG	0.96	0.97	0.97	
	NET	1.00	0.96	0.96	

Multi Nomial Naïve Bayes	Kelas	Precision	Recall	F1 Score	Accuracy
	POS	0.82	0.87	0.85	
	NEG	0.60	0.97	0.74	
	NET	1.00	0.30	0.46	

Dengan melakukan *hyperparameter optimization* pada SVM didapatkan bahwa terdapat hasil peningkatan yang signifikan dari performa model SVM. Didapatkan bahwa SVM dapat mengungguli Multinomial Naïve Bayes dengan selisih akurasi sebesar 25% dimana sebelumnya dengan metode evaluasi train test split Multinomial Naïve Bayes dapat memberikan performa yang lebih baik disbanding SVM tanpa optimalisasi *hyperparameter*. Hal ini menunjukkan bahwa performa SVM jauh lebih baik setelah dilakukan *hyperparameter tuning*. Dengan Precision dari setiap kelas diatas 0.9 dan akurasi 97% dapat dikatakan bahwa algoritma SVM dapat melakukan klasifikasi terhadap sentiment masyarakat dengan baik.

Tabel 12. Hasil Skenario Evaluasi SVM dengan *parameter tuning* dan *Multinomial Naïve Bayes* dengan *cross validation*

SVM	Index K	Accuracy	Precision	Recall	F1 Score
	1	0.76	0.85	0.77	0.76
	2	0.99	0.99	0.99	0.99
	3	0.97	0.97	0.97	0.97
	4	0.98	0.98	0.98	0.98
	5	1.00	1.00	1.00	1.00
	6	0.94	0.95	0.94	0.94
	7	0.96	0.97	0.96	0.96
	8	1.00	1.00	1.00	1.00
	9	0.99	0.99	0.99	0.99
	10	0.93	0.94	0.94	0.94
rata-rata		0.95	0.96	0.95	0.95

Multi Nomial Naïve Bayes	Index K	Accuracy	Precision	Recall	F1 Score
	1	0.53	0.67	0.54	0.45
	2	0.72	0.80	0.72	0.65
	3	0.68	0.77	0.68	0.63
	4	0.71	0.80	0.71	0.64
	5	0.74	0.81	0.74	0.69
	6	0.69	0.80	0.68	0.63
	7	0.68	0.78	0.68	0.60
	8	0.67	0.77	0.67	0.56
	9	0.66	0.77	0.66	0.58
	10	0.67	0.77	0.68	0.61
rata-rata		0.67	0.77	0.68	0.60

Pada evaluasi cross validation dan menggunakan mean approach didapatkan bahwa performa dari SVM meningkat dengan signifikan dan memiliki nilai yang lebih baik dibanding dengan Multinomial Naïve bayes dari segi akurasi, recall dan F-1 Score. Nilai maximum akurasi, precision recall dan F1 Score secara berurutan dari SVM adalah 1.0 , 1.0 , 1.0 dan 1.0 sedangkan dengan naïve bayes secara berurutan adalah 0.74, 0.81, 0.74, 0.69 . Jika menggunakan Optimistic approach, dimana nilai maximum dari setiap matrix pada index K yang diambil maka

SVM dapat dikatakan dapat mengungguli Multinomial Naïve bayes dengan nilai sempurna . dengan rata rata hasil akurasi sebesar 0.95 , precision sebesar 0.96 , recall sebesar 0.95 dan F-1 Score sebesar 0.95 dapat dikatakan bahwa SVM dapat melakukan klasifikasi terhadap sentiment masyarakat Indonesia mengenai RKUHP

4.2 Analisis Hasil Pengujian

Berdasarkan penelitian yang dilakukan, SVM yang dibangun tanpa menggunakan *hyperparameter* tuning belum memberikan performa yang baik jika dibandingkan dengan algoritma *multinomial naïve bayes* . Dengan melakukan *grid search* terhadap SVM didapatkan parameter terbaik nilai $C=1000$, $\gamma=0.0001$, kernel=*Radial Basis Function* . dengan melakukan *hyperparameter tuning* dan melatih ulang model SVM , didapatkan hasil yang signifikan dengan meningkatkan akurasi, *precision*, *recall* dan *f1 score* dengan selisih akurasi 25% , *precision* sebanyak 24% , *recall* sebanyak 25% dan *f1 score* sebanyak 25% . dengan rata rata hasil akurasi sebesar 0.95 , precision sebesar 0.96 , recall sebesar 0.95 dan F-1 Score sebesar 0.95 dapat dikatakan bahwa SVM dapat melakukan klasifikasi terhadap sentiment masyarakat Indonesia mengenai RKUHP

5. Kesimpulan

Berdasarkan hasil penelitian yang dilakukan maka didapatkan kesimpulan sebagai berikut :

1. Dalam pengujian ini didapatkan bahwa performa dari algoritma SVM akan meningkat dengan signifikan dengan melakukan *hyperparameter tuning* menggunakan *grid search*
2. *Hyperparameter* dari SVM yang mempengaruhi performa dari SVM pada pengujian kali ini adalah C , γ dan kernel yang digunakan. Didapatkan nilai kombinasi *hyperparameter* terbaik yaitu $C=1000$, $\gamma=0.0001$, kernel= *Radial Basis Function*
3. Dapat disimpulkan bahwa SVM yang dibangun dengan menggunakan kombinasi *hyperparameter* terbaik dapat mengukur sentimen dari masyarakat terhadap RKUHP dengan hasil akurasi sebesar 95%
4. dengan rata rata hasil akurasi sebesar 0.95 , precision sebesar 0.96 , recall sebesar 0.95 dan F-1 Score sebesar 0.95 dapat dikatakan bahwa SVM dapat melakukan klasifikasi terhadap sentiment masyarakat Indonesia mengenai RKUHP

Dengan berhasilnya model SVM melakukan klasifikasi terhadap sentiment masyarakat Indonesia mengenai RKUHP , didapatkan hasil bahwa sentiment dari masyarakat Indonesia terhadap RKUHP cenderung diidentifikasi menjadi dua kelas positif dan negative. Jika dilihat dari kuantitas label yang diberikan oleh annotator diidentifikasi bahwa sentiment masyarakat Indonesia cenderung Negatif terhadap RKUHP. Dan Precision yang didapatkan oleh kelas negative terbilang berada di urutan tertinggi kedua setelah kelas Netral, sehingga dapat dikatakan bahwa model cenderung mengidentifikasi kelas netral dan negative. Berdasarkan jumlah data label dan data hasil prediksi dari model dengan precision dan akurasi yang tinggi dapat dikatakan bahwa masyarakat cenderung memberikan sentiment negative terhadap RKUHP

Untuk penelitian yang akan datang, model yang dibuat akan lebih baik jika terdapat nilai kesepakatan (*kappa score*) antara pelabel yang melakukan pelabelan.

Referensi

- [1] J. Vassileva, "Motivating participation in social computing applications: A user modeling perspective," *User Model. User-adapt. Interact.*, vol. 22, no. 1–2, pp. 177–201, 2012.
- [2] Y. Yang, Y. Yang, B. J. Jansen, and M. Lalmas, "Computational advertising: A paradigm shift for advertising and marketing?," *IEEE Intell. Syst.*, vol. 32, no. 3, pp. 3–6, 2017.
- [3] C. Eds, *Twitter Book*. 2014.
- [4] B. Liu, *Sentiment Analysis: A Fascinating Problem*. 2012.
- [5] L. Arliman S, "Kodifikasi RUU KUHP Melemahkan Komisi Pemberantasan Korupsi," *Uir Law Rev.*, vol. 2, no. 01, p. 256, 2018.
- [6] "Keywords 2 . RELATED WORK," pp. 1–7, 2012.
- [7] F. Ugm and F. Ugm, "Analisis Sentimen Twitter untuk Teks Berbahasa Indonesia dengan Maximum Entropy dan Support Vector Machine," *IJCCS (Indonesian J. Comput. Cybern. Syst.*, vol. 8, no. 1, pp. 91–100, 2014.
- [8] R. Varghese and M. Jayasree, "Aspect based Sentiment Analysis using support vector machine classifier," *Proc. 2013 Int. Conf. Adv. Comput. Commun. Informatics, ICACCI 2013*, pp. 1581–1586, 2013.
- [9] W. S. Noble, "What is a support vector machine?," *Nat. Biotechnol.*, vol. 24, no. 12, pp. 1565–1567, 2006.
- [10] E. Osuna, R. Freund, and F. Girosi, "An improved training algorithm for support vector machines BT - Neural Networks for Signal Processing VII," pp. 276–285, 1997.
- [16] K. Sigit, A. P. Dewi, G. Windu, Nurmalasari, T. Muhamad, and N. Kadinar, "Comparison of Classification Methods on Sentiment Analysis of Political Figure Electability Based on Public Comments on Online News Media Sites," *IOP Conf. Ser. Mater. Sci. Eng.*, vol. 662, no. 4, 2019, doi: 10.1088/1757-899X/662/4/042003.
- [17] A. Krouska, C. Troussas, and M. Virvou, "The effect of preprocessing techniques on Twitter sentiment analysis," *IISA 2016 - 7th Int. Conf. Information, Intell. Syst. Appl.*, no. November 2017, 2016, doi: 10.1109/IISA.2016.7785373.
- [11] M. Hasnain, M. F. Pasha, I. Ghani, M. Imran, M. Y. Alzahrani, and R. Budiarto, "Evaluating Trust Prediction and Confusion Matrix Measures for Web Services Ranking," *IEEE Access*, vol. 8, pp. 90847–90861, 2020, doi: 10.1109/ACCESS.2020.2994222.
- [13] M. Ramya and J. A. Pinakas, "Different Type of Feature Selection for Text Classification," *Int. J. Comput. Trends Technol.*, vol. 10, no. 2, pp. 102–107, 2014, doi: 10.14445/22312803/ijctt-v10p118.
- [14] C. Vania, M. Ibrahim, and M. Adriani, "Sentiment Lexicon Generation for an Under-Resourced Language," *Int. J. Comput. ...*, vol. 5, no. 1, pp. 59–72, 2014.
- [15] F. Koto and G. Y. Rahmaningtyas, "Inset lexicon: Evaluation of a word list for Indonesian sentiment analysis in microblogs," *Proc. 2017 Int. Conf. Asian Lang. Process. IALP 2017*, vol. 2018-January, pp. 391–394, 2018, doi: 10.1109/IALP.2017.8300625.
- [18] I. Syarif, A. Prugel-Bennett, and G. Wills, "SVM parameter optimization using grid search and genetic algorithm to improve classification performance," *Telkomnika (Telecommunication Comput. Electron. Control.*, vol. 14, no. 4, pp. 1502–1509, 2016, doi: 10.12928/TELKOMNIKA.v14i4.3956.
- [12] Rossi ALD, de Carvalho ACP. Bio-inspired Optimization Techniques for SVM Parameter Tuning. In *10th Brazilian Symposium on Neural Networks*, 2008. SBRN '08. Presented at the 10th Brazilian Symposium on Neural Networks, 2008. SBRN '08. 2008: 57-62.

Lampiran

