

Analisis Sentimen Terhadap Ulasan Film Menggunakan Word2Vec dan SVM

Yusuf Surya T¹, Said Al Faraby², Mahendra Dwifabri³

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

yusufst@students.telkomuniversity.ac.id¹, saidalfaraby@telkomuniversity.ac.id²,
mahendradp@telkomuniversity.ac.id³

Abstrak

Pertumbuhan dan penyebaran informasi pada saat ini didukung dengan teknologi yang semakin berkembang. Informasi dapat menyebar luas di internet dengan cepat. Salah satunya informasi opini tentang sebuah film. Ada yang beropini positif maupun negatif. Terdapat polaritas dalam suatu movie review dikarenakan setiap orang mempunyai opininya masing-masing. Akibatnya banyak penikmat movie yang kesulitan menemukan informasi yang sesuai dengan kebutuhannya. Dengan adanya permasalahan tersebut maka metode yang tepat untuk menganalisisnya adalah analisis sentimen. Dalam penelitian ini, dataset analisis sentiment kemudian dilakukan tahap *preprocessing*, ekstraksi fitur ekstraksi fitur Word2vec dan kemudian dilakukan klasifikasi menggunakan metode Support Vector Machine. Dengan sistem yang dibangun maka menghasilkan nilai akurasi terbaik sebesar 78.75% dan F1-score terbaik sebesar 78,74%. Sistem yang dibangun menggunakan *lemmatization* dengan 300 dimensi Word2vec, beserta klasifikasi SVM linier memiliki nilai performansi tertinggi.

Kata kunci : analisis sentimen, Word2Vec, SVM

Abstract

The growth and dissemination of information at this time is supported by increasingly developing technology. Information can spread widely on the internet quickly. One of them is opinion information about a film. There are positive and negative opinions. There is a polarity in a movie review because everyone has their own opinion. As a result, many movie lovers have difficulty finding information that suits their needs. With these problems, the right method to analyze it is sentiment analysis. In this study, the sentiment analysis dataset was then carried out in the preprocessing stage, feature extraction of Word2vec feature extraction and then classification using the Support Vector Machine method. With the system built, it produces the best accuracy value of 78.8% and the best F1-score of 78.79%. The system built using lemmatization with 300 Word2vec dimensions, along with the linear SVM classification has the highest performance value

Keywords: sentiment analysis, Word2Vec, SVM

1. Pendahuluan

Latar Belakang

Pesatnya perkembangan teknologi mempengaruhi pertumbuhan dan penyebaran informasi. Teknologi internet berperan penting dalam penyebaran informasi. Informasi teks yang terdapat pada web dapat dibagi menjadi dua kategori yaitu fakta atau opini. Fakta merupakan informasi yang berdasarkan objektivitas sementara opini merupakan ekspresi *sentiment* dari penulisnya[3].

Pendapat yang diungkapkan di internet akan menjadi informasi yang penting, salah satunya informasi ulasan suatu film. Pengguna internet dapat dengan mudah menemukan informasi tertulis tentang sebuah film sebagai referensi film untuk ditonton. Dalam satu bulan, 69% penonton dapat menonton satu film dan 17,5% penonton dapat menontonnya dua kali film yang sama [1]. Informasi tentang ulasan film tersebar di internet. Akibatnya, banyak pecinta film kesulitan mencari informasi yang mereka butuhkan. Karena setiap orang memiliki opininya masing-masing, informasi yang didistribusikan cenderung tidak merata. Berdasarkan permasalahan yang dihadapi oleh para penikmat film maka penelitian analisis sentiment sangat akan mempermudah pengguna. Analisis sentiment adalah jenis pemrosesan bahasa alami untuk melacak mood publik tentang produk atau topik tertentu [2]. Tujuannya untuk menentukan polaritas dalam suatu movie review berdasarkan data yang membicarakan tentang suatu movie. Teknik analisis sentimen digunakan untuk mengekstrak dan mengelola data kemudian memasukkannya ke dalam kategori negatif atau positif. Analisis sentimen termasuk golongan supervised learning. Data movie review yang sudah dikumpulkan dijadikan data training dan data testing. Selanjutnya dilakukan proses klasifikasi.

Untuk mengetahui secara langsung apa maksud dari sebuah kalimat dalam hasil review, diperlukan algoritma yang akan memproses dan memetakan informasi tersebut menjadi sebuah vektor yaitu dengan menggunakan metode Word2vec. Pada penelitian sebelumnya Word2Vec banyak digunakan untuk melakukan klasifikasi berdasarkan vektor dari data teks [4]. Sementara pada penelitian ini Word2Vec digunakan untuk melihat vektor dari sebuah paragraf atau dokumen yang kemudian akan diklasifikasi dengan menggunakan metode *Support*

Vector Machine (SVM). Dalam penelitian ini, metode SVM digunakan untuk menganalisa sentimen dari hasil review dalam forum website IMDB. *Sentiment Analysis* juga merupakan bidang ilmiah yang dapat menganalisis opini, sentimen, evaluasi, penelitian, penilaian, sikap, dan emosi orang tentang entitas seperti produk, layanan, organisasi, individu, masalah, peristiwa, topik, dan atribut lainnya[5]. Terdapat tiga level dalam *sentiment analysis* yaitu *coarse grained sentiment analysis*, *fine-grained sentiment analysis* dan *adjective sentiment analysis* [6]. Penulis memilih metode Word2Vec untuk melakukan klasifikasi, karena metode ini dapat membantu komputer untuk mengidentifikasi kombinasi kata yang akan diklasifikasikan. Pada penelitian sebelumnya[4], penggunaan Word2Vec memiliki hasil akurasi yang baik. Proses klasifikasi dilakukan dengan menggunakan metode SVM karena adanya penelitian sebelumnya dengan menggunakan metode SVM dengan hasil akurasi yang cukup baik yaitu 85.6%[7].

Topik dan Batasannya

Topik dan batasan masalah dalam penelitian tugas akhir ini adalah mengetahui hasil pengujian performansi dari model berdasarkan skenario dalam penelitian ini. Batasan masalah pada penelitian ini adalah dilakukan analisis sentimen dengan menggunakan Word2Vec dengan metode SVM. Data yang digunakan adalah data review film berbahasa Inggris yang berasal dari website IMDB, dengan jumlah 10.000 data ulasan dengan label positif dan negatif.

Tujuan

Tujuan dalam penelitian ini adalah membangun sistem analisis sentimen pada data ulasan film berbahasa Inggris menggunakan Word2Vec dengan metode SVM. Serta, mengetahui performansi dari sistem penelitian yang dibangun.

Organisasi Tulisan

Pada bagian selanjutnya pada bab 2 akan membahas studi literatur yang berkaitan dengan penelitian ini. Pada bab 3 akan membahas teori dan rancangan alur sistem penelitian. Dan pada bab 4 dan bab 5 akan membahas hasil dan analisis pada penelitian ini dan juga kesimpulan penelitian.

2. Studi Terkait

Pada Penelitian ini, terdapat beberapa referensi mengenai penelitian sebelumnya terkait yang memiliki fokus terhadap Sentiment Analysis, Movie Review, ataupun metode yang diusulkan pada penelitian ini. Adapun beberapa penelitian terkait [3,15,8,14] yaitu:

Penelitian analisis sentimen sudah pernah dilakukan oleh Juliansyah et al[15] menggunakan metode KNN dengan *negation handling* pada dataset *movie review*. Penelitian ini menggunakan ekstraksi fitur TF-IDF, yang selanjutnya dilakukan klasifikasi menggunakan metode KNN. Tujuan penelitian ini untuk mengetahui performansi sistem yang dibangun dan pengaruh metode *negation handling* terhadap sistem yang dibangun menggunakan dataset *review movie*. Penelitian ini menghasilkan nilai akurasi terbaik sebesar 53.6% dan nilai *F1-score* sebesar 53.1%. Kemudian sistem yang dibangun tanpa menggunakan metode *negation handling* namun menggunakan metode *stemming* dan *mutual information* adalah sistem yang memiliki nilai performansi terbaik.

Penelitian yang mengimplementasikan SVM dilakukan oleh Winda et al[3] menggunakan dataset *review movie* sebanyak 2000 ulasan dalam berbahasa Inggris yang dilabeli positif dan negatif secara manual. Tujuan dari penelitian tersebut adalah menganalisis komposisi data *Training* dan data *Testing*, juga menganalisis fungsi *Kernel* SVM. Hasil dari penelitian ini, menyimpulkan bahwa Komposisi data *Training* mempengaruhi akurasi, hal ini disebabkan semakin tinggi komposisi data *training* maka jumlah variasi data yang di train akan semakin banyak sehingga, sistem dapat melakukan klasifikasi lebih baik. Pada sistem ini nilai F1-Measure terbaik pada komposisi data *training* 800 dan data *testing* 200. Dan hasil dari analisis fungsi *Kernel* pada SVM yang dapat digunakan pada kasus *review film* yaitu *kernel linear*, *kernel RBF* dan *kernel polynomial*. Dengan hasil akurasi tertinggi pada dataset ini yaitu pada *kernel polynomial* dengan akurasi 54.19%

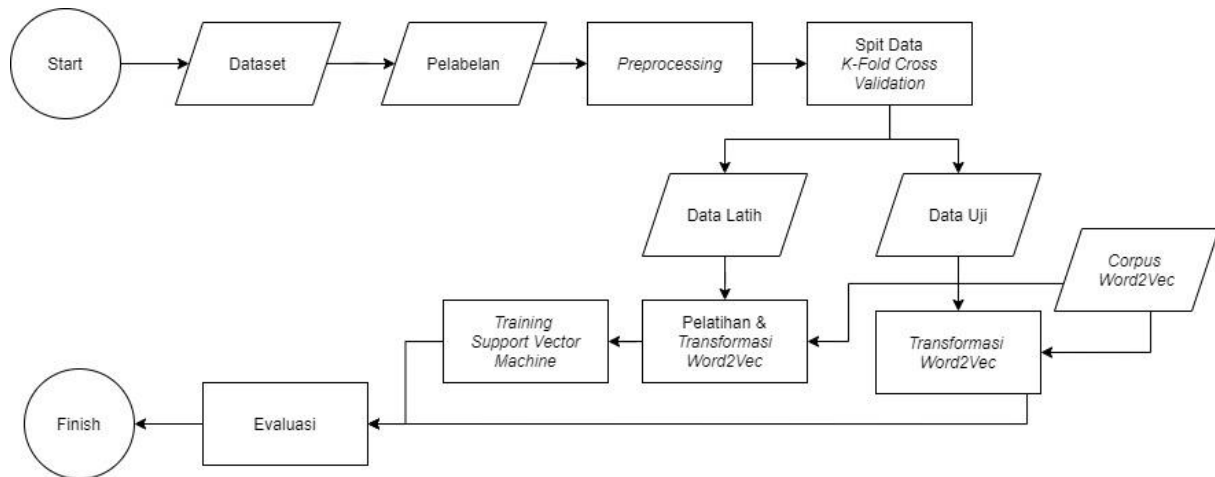
Penelitian mengenai *sentiment analysis* menggunakan model Word2vec dilakukan oleh Fauzi et al[8]. Namun pada penelitiannya penggunaan model word2vec menghasilkan akurasi 70%, karena data yang digunakan sedikit. Dalam data yang sedikit word2vec tidak dapat menangkap kemiripan makna dengan baik. Sehingga dilakukan penelitian terkait yang mana menggunakan data Facebook dan juga data artikel Wikipedia berbahasa Indonesia sebagai model word2vec. Pada penelitian ini dilakukan perbandingan ekstraksi fitur *Average base Word2vec* dan *Bag of Centroid* base word2vec dan juga dilakukan penggabungan keduanya kemudian dilakukan klasifikasi menggunakan metode Support Vector Machine.

Penelitian yang mengimplematisasikan *Information Gain* dan *Naive Bayes* dilakukan oleh Laila et al[14] menggunakan dataset ulasan buku berbahasa Inggris sebanyak 1000 data yang dilabeli positif dan negatif. Tujuan penelitian ini mengetahui performa yang dihasilkan dari penerapan klasifikasi sentimen menggunakan *Information Gain* dan *Naive bayes*. Hasil penelitian pada tahap *preprocessing* tanpa menggunakan *lemmatization* menghasilkan nilai F1-score sebesar 86.77%, sedangkan menggunakan *preprocessing lemmatization* menghasilkan akurasi

sebesar 88%. Hal ini disebabkan *lemmatization* mengubah kata menjadi kata dasar. Dengan menggunakan *lemmatization*, sistem akan dapat mengenali fitur yang sama, sehingga tidak ada kata yang sebenarnya memiliki bentuk dasar yang sama namun dianggap berbeda.

3. Sistem yang Dibangun

Penelitian ini akan dibangun model analisis sentimen dengan menggunakan ekstraksi fitur Word2Vec Skip-gram dan klasifikasi menggunakan *Support Vector Machine*(SVM). Alur penelitian ini diilustrasikan pada Gambar 1:



Gambar 1. Sistem pada Analisis Sentimen pada Ulasan Film menggunakan Word2Vec dan SVM

3.1. Dataset

Data yang akan digunakan pada penelitian ini adalah hasil kumpulan *review film* yang didapatkan dari *website forum ulasan film internasional* yaitu *imdb.com* yang telah dikumpulkan oleh Andrew Mass. Dataset yang digunakan sebanyak 10.000 ulasan, yang meliputi data ulasan positif sebanyak 5028 ulasan, dan data ulasan negatif sebanyak 4972 ulasan. Berikut contoh dataset yang digunakan dalam penelitian ini yang dijabarkan pada tabel 1:

Tabel 1. Dataset Penelitian

Label	Ulasan
Positif	<i>"though it is by no means his best work, laissez-passer is a distinguished and distinctive effort by a bona-fide master , fascinating film replete with reward to be had by all willing to make the effort to reap them "</i>
Negatif	<i>"a visually flashy but narratively opaque and emotionally vapid exercise in style and mystification . the story is also as unoriginal as they come , already having been recycled more times than i'd care to count . "</i>

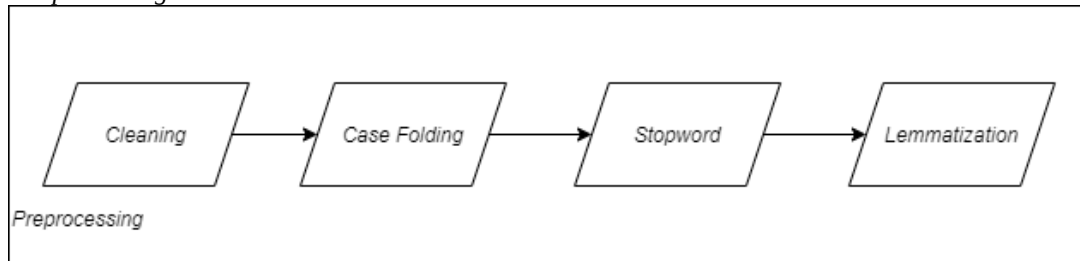
3.2. Pelabelan

Pada penelitian ini, ada 2 jenis pelabelan yang akan digunakan yaitu positif, dan negatif. Dalam satu dokumen ini akan diidentifikasi apakah ulasan film tersebut positif atau negatif, jika positif akan diberi label "1" dan jika negatif diberi label "0". Berikut contoh pelabelan dari ulasan film yang digunakan dijabarkan pada tabel 2:

Tabel 2. Contoh Pelabelan Ulasan

Ulasan	Label
<i>"though it is by no means his best work, laissez-passer is a distinguished and distinctive effort by a bona-fide master , fascinating film replete with reward to be had by all willing to make the effort to reap them "</i>	1
<i>"a visually flashy but narratively opaque and emotionally vapid exercise in style and mystification . the story is also as unoriginal as they come , already having been recycled more times than i'd care to count . "</i>	0

3.3. Preprocessing



Gambar 2. Tahap Preprocessing

Setelah melewati tahap pelabelan data, rancangan sistem dalam tugas akhir ini yaitu tahap preprocessing. Tahap ini dilakukan untuk mengolah dataset mentah yang belum terstruktur dan belum siap melewati tahapan klasifikasi. Ditahap ini data akan dibersihkan dan dibentuk menjadi lebih terstruktur sehingga siap untuk melewati tahap klasifikasi. Tahapan dalam preprocessing yang dilakukan dalam penelitian ini yaitu cleaning, case folding, stopwords, dan lemmatization. Tahapan preprocessing dalam penelitian ini dapat dilihat pada gambar 2.

3.3.1 Cleaning

Cleaning merupakan proses untuk menghilangkan angka, double spasi, tanda baca, dan simbol yang tidak akan dibutuhkan dalam penelitian. Contoh pada kalimat “If you like original gut wrenching laughter, You will like this MOVIE!!!!.” menjadi “If you like original gut wrenching laughter You will like this MOVIE”.

3.3.2 Case folding

Case folding merupakan proses mengubah semua kata pada dokumen dataset menjadi *lowercase* atau huruf kecil sehingga sudah tidak huruf kapital. Misalnya pada kalimat “If you like original gut wrenching laughter You will like this MOVIE” menjadi “if you like original gut wrenching laughter you will like this movie”.

3.3.3 Stopword

Stopword merupakan proses menghilangkan kata yang tidak memiliki arti. Pada proses *stopword* digunakan kamus *stopword* bahasa Inggris. Contoh penggunaan *stopword* pada kalimat “if you like original gut wrenching laughter you will like this movie” akan menjadi “like original gut wrenching laughter like movie”.

3.3.4 Lemmatization

Lemmatization merupakan proses mengubah kata menjadi kata dasar dengan cara menghapus imbuhan atau sisipan kata [9]. Pada proses *lemmatization* ini menggunakan kamus *WordNetLemmatizer*, karena dataset dalam penelitian ini menggunakan dataset berbahasa Inggris. Contoh penggunaan *lemmatization* pada kalimat “like original gut wrenching laughter like movie” akan menjadi “like original gut wrench laughter like movie”.

Tabel 3. Contoh potongan data melewati tahap preprocessing

Dataset	If you like original gut wrenching laughter, you will like this movie. If you are young or old then you will love this movie!!!.
Cleaning & Case folding	if you like original gut wrenching laughter you will like this movie if you are young or old then you will love this movie
Stop Word	like original gut wrenching laughter like movie young old love movie
Lemmatization	like original gut wrench laughter like movie young old love movie

3.4. Data Split

Setelah dataset melakukan tahap preprocessing, Dataset akan dipisahkan menjadi dua kategori yaitu data latih dan data uji, dalam penelitian ini split data akan menggunakan *k-fold* dengan perbandingan 8:2, dengan detail data latih sebanyak 8000 data ulasan film dengan label positif, dan label negatif dan data uji sebanyak 2000 data ulasan film. Berikut persebaran antara data latih dan data uji setelah data dipisahkan pada tabel :

Tabel 4. Persebaran Data Latih dan Data Uji

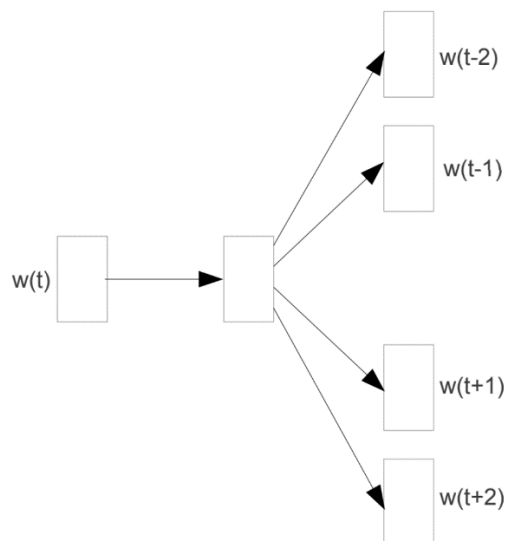
Kategori	Label	
	Negatif	Positif
Data Latih	3997	4003
Data Uji	975	1025

3.5. Word2Vec

Pada tahap selanjutnya data akan dilakukan ekstraksi fitur. Dalam penelitian ini metode yang digunakan ialah Word2Vec Skip-gram *negative* sampling. Secara definisi Word2Vec adalah algoritma untuk mengubah kata menjadi vektor menggunakan *neural network* untuk mempelajari *embedding* kata [11]. Diubah nya teks menjadi vektor supaya mesin bisa lebih memahami makna setiap kata lebih akurat. Sehingga kata-kata yang memiliki arti yang cukup mirip akan memiliki output yang berdekatan antara satu sama lain. *Input* dan *output* dalam proses Word2Vec ini ialah kumpulan vektor. Word2Vec memiliki dua model arsitektur yaitu *Continuous Bag Of Word* (CBOW) dan Skip-gram. Proses arsitektur model Skip-gram yaitu dengan membuat prediksi kata-kata di sekitar pada suatu kalimat, sedangkan arsitektur model CBOW memprediksi kata-kata yang berada dalam satu kalimat[9], sehingga Skip-gram merupakan metode yang lebih efisien untuk mempelajari representasi vektor kata dalam jumlah besar pada teks yang tidak terstruktur. Dalam penelitian ini Word2Vec yang digunakan ialah model Skip-gram dikarenakan dapat memaksimalkan rata-rata *log probability*. Berikut persamaan dari model Skip-gram:

$$\frac{1}{T} \sum_{t=1}^T \sum_{-c \leq j \leq c, j \neq 0} \log p(w_{t+j} | w_t) \quad (1)$$

Dengan w_t merupakan kata *center*, w_{t+j} merupakan kata setelah *center*, dan c adalah ukuran *training context*. Implementasi gambar arsitektur *skip-gram* dari hasil persamaan berikut berada pada gambar 3.

**Gambar 3. Arsitektur Skip-gram**

Dalam penelitian ini model Word2vec akan berfokus pada model *Skip-gram Negative Sampling* (SGNS) 100 dimensi dan 300 dimensi *corpus* yang berasal dari *corpus Wikipedia*[12]. SGNS sendiri merupakan model yang berguna untuk meningkatkan kecepatan serta kualitas dari komputasi Word2Vec[11].

3.6. Klasifikasi Support Vector Machine

SVM merupakan *supervised learning*. SVM pertama kali diperkenalkan oleh Vapnik. SVM bertujuan untuk mencari *hyperplane* optimal yang berperan sebagai pemisah kelas pada ruang input dengan cara memaksimalkan jarak antar kelas dengan mencari margin, atau support vector yang merupakan titik data yang dekat dengan *hyperplane* tersebut[13]. Upaya mencari lokasi *hyperplane*

optimal ini usaha yang digunakan untuk menemukan atau mencari letak atau lokasi hyperplane ini merupakan inti dari proses pada *Support Vector Machine* Berikut formulasi untuk mengoptimalkan nilai margin pada SVM menggunakan persamaan 2:

$$m = \min_w \frac{1}{2} (\|w\|)^2 \quad (2)$$

Dimana:

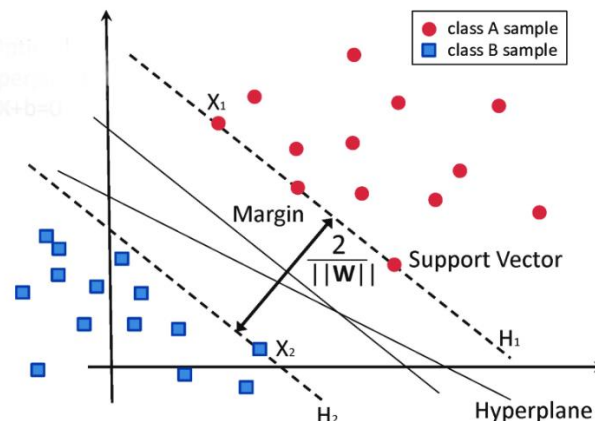
$$y_i(x_i * w + b) - 1 \geq 1, i = 1, \dots, l, \quad (3)$$

&

$$w * x + b \leq -1 \quad (4)$$

Dimana x_i merupakan data masukan, y_i merupakan data keluaran, sedangkan w dan b merupakan parameter yang akan dicari nilainya.

Support Vector Machine digunakan untuk menyelesaikan masalah klasifikasi linear dan non linear. SVM secara non linear menggunakan kernel trick yang berfungsi untuk membantu dan memudahkan dalam melakukan klasifikasi dalam bentuk non linear. Kernel digunakan untuk memetakan inputan di *input space* ke dalam *feature space* membuat pemisah antar kelas bukan lagi garis lurus merupakan garis lengkung atau bidang dalam dimensi yang lebih tinggi. Tentu pendekatan dengan kernel ini berbeda dengan metode lain karena biasanya pendekatan dengan metode lain mengurangi dimensi awal untuk menyederhanakan proses komputasi.



Gambar 4. Ilustrasi *hyperplane* sebagai pemisah dalam dua kelas

Dalam penelitian ini metode SVM yang digunakan ialah SVM linear dan SVM non linear yaitu kernel RBF dan kernel Sigmoid. Berikut adalah beberapa fungsi kernel yang digunakan, yaitu :

Kernel Linear : $K(x_i, x) = x_i^T x$ (5)

Kernel *Radial Basis Function* : $K(x_i, x) = \exp(-\gamma |x_i - x|^2)$, (6)

Kernel Sigmoid : $K(x_i, x) = \tanh(\gamma x_i^T x + r)$ (7)

3.7. Evaluasi

Dasar penelitian ini adalah proses bentuk klasifikasi, sehingga hasil dari proses klasifikasi yang telah dilakukan harus dievaluasi. Tahap evaluasi merupakan tahap terakhir serangkaian tahap dari perancangan sistem, dalam penelitian ini penulis menggunakan metode evaluasi *k-fold (5-fold) cross validation* serta *confusion matrix*. *Confusion matrix* digunakan untuk mencari nilai *accuracy*, *precision*, *recall*, dan *f1-score*. Sedangkan *k-fold cross validation* digunakan untuk memastikan ulang nilai yang didapat dan dicari nilai rata-ratanya. Nilai hasil klasifikasi yang digunakan untuk menentukan *accuracy*, *precision*, *recall*, dan *f1-score* pada persamaan berikut:

$$accuracy = \frac{TP+TN}{TP+TN+FP+FN} \quad (8)$$

$$precision = \frac{TP}{TP+FP} \quad (9)$$

$$recall = \frac{TP}{TP+FN} \quad (10)$$

$$f1-Score = 2 \times \frac{(Presisi \times Recall)}{(Presisi+Recall)} \quad (11)$$

Nilai *True Positives* (TP), *True Negatives* (TN), *False Positives* (FP), dan *False Negatives* (FN) dibutuhkan dari *confussion matrix* untuk menghitung point-point yang telah disebutkan untuk tahap evaluasi.

4. Evaluasi

Penelitian ini dilakukan dalam beberapa tahap, yaitu pertama proses *preprocessing* pada dataset. Setelah data melewati tahap *preprocessing* data akan masuk tahap berikutnya yaitu data split dengan persebaran data latih sebesar 80% dan data uji sebesar 20%, lalu data latih dan data uji akan melakukan proses ekstraksi fitur word2vec skip-gram *negative sampling* 100 dimensi dan 300 dimesni. Selanjutnya data akan masuk tahap selanjutnya yaitu pembangunan model klasifikasi menggunakan metode *Support Machine Vector* linear dan non linear dan divalidasi pada tahap evaluasi menggunakan *k-fold 5-fold cross validation* dan *confussion matrix*. Terdapat beberapa skenario pengujian yang dilakukan, yang dijabarkan seperti berikut:

- Skenario 1 : Pengujian peromansi tahap *Preprocessing Lemmatization* dan tanpa *Preprocessing Lemmatization*.
- Skenario 2 : Pengujian peromansi penggunaan dimensi 100 dan dimesni 300 pada Word2Vec.
- Skenario 3 : Pengujian peromansi penggunaan SVM lienar, kernel RBF dan kernel sigmoid.

4.1 Skenario 1 Pengujian Tahap *Preprocessing Lemmatization*

Pada skenario 1 dilakukan pengujian untuk membandingkan tahap *preprocessing* menggunakan *lemmatization* dan tanpa *lemmatization*. Hasil dari skenario ini adalah dataset dengan *lemmatization* dan dataset tanpa *lemmatization*.

Tabel 5. Hasil Pengujian Tahap *Preprocessing*

<i>Preprocessing</i>	Akurasi	<i>Precision</i>	<i>Recall</i>	F1- Score	F1-Score Max
<i>Lemmatization</i>	78.75%	78.76%	78.77%	78.74%	78.74%
No <i>Lemmatization</i>	77.69%	77.71%	77.71%	77.66%	

Berdasarkan hasil dari skenario 1, hasil pengujian dataset menggunakan *lemmatization* memiliki nilai akurasi, precision, recall dan f1-score lebih tinggi dibanding dataset yang tidak menggunakan *lemmatization*. Hal ini disebabkan pada data yang melewati tahap *preprocessing* yang menggunakan *lemmatization* melakukan penghapusan imbuhan kata yang tidak dibutuhkan sesuai *corpus*. Contoh pemotongan fitur yang dilakukan pada *lemmatization* yaitu kata 'wear it' dan 'worn' menjadi 'wear', sedangkan pada tahap *preprocessing* tanpa *lemmatization* apabila terdapat kata pada 'wear it' dan 'worn' tidak dilakukan nya pemotngan imbuhan kata walalu memeilki arti yang cukup berbeda. Sehingga menyebabkan nilai akurasi dan f1-score menurun.

4.2 Skenario 2 Pengujian Penggunaan Dimensi 100 dan Dimensi 300 pada Word2Vec

Hasil dari skenario 1 yaitu dataset yang menggunakan *preprocessing lemmatization* akan digunakan dalam skenario 2. Dalam skenario 2 ini dilakukan pengujian penggunaan dimensi word2vec antara 100 dimensi dengan 300 dimensi.

Tabel 6. Hasil Pengujian Penggunaan Dimensi pada Word2Vec

Word2Vec	Akurasi	<i>Precision</i>	<i>Recall</i>	F1- Score	F1-Score Max
100 Dimensi	74.63%	74.6%	74.64%	74.62%	78.74%
300 Dimensi	78.75%	78.76%	78.77%	78.74%	

Berdasarkan hasil pengujian skenario 2, penggunaa dimensi word2vec menggunakan 100 dimensi memiliki hasil lebih rendah dibanding dengan penggunaan 300 dimensi yang dimaan memiliki selisih nilai f1-score sebesar 4.12%. Hal ini disebabkan Hal ini disebabkan pada penggunaan dimensi 300 nilai *recall* yang bekerja untuk memprediksi data kelas aktual dan dan prediksinya positif dan nilai *precision* yang bekerja untuk memprediksi benar positif dibandingkan dengan keseluruhan hasil yang diprediksi

positif lebih tinggi. Sehingga dalam penelitian ini menyimpulkan penggunaan 300 dimensi jelas meningkatkan nilai performansi pada model word2vec skip-gram *negative sampling*.

4.3 Skenario 3 Pengujian Perbandingan Penggunaan Model SVM Linier dan Non Linier

Hasil dari skenario 1 dan skenario 2 yaitu dataset yang menggunakan *preprocessing lemmatization* dan word2vec skip-gram *negative sampling* 300 dimensi akan digunakan dalam skenario 3. Dalam skenario 3 akan dilakukan pengujian antara penggunaan model SVM linear, kernel RBF, dan kernel sigmoid.

Tabel 7. Hasil Pengujian Penggunaan Model SVM Linear dan Non Linear

Klasifikasi SVM	Akurasi	Precision	Recall	F1- Score	F1-Score Max
Linear SVM	78.75%	78.76%	78.77%	78.74%	78.74%
Kernel RBF	77.78%	77.78%	77.8%	77.77%	
Kernel Sigmoid	74.89%	74.91%	74.92%	74.89%	

Berdasarkan hasil pengujian skenario 3, hasil dari penggunaan metode SVM linear memiliki hasil performansi lebih tinggi dibanding dengan kernel RBF dan kernel sigmoid yang dimana menghasilkan nilai f1-score sebesar 78.74% untuk SVM linear, 77.77% untuk SVM kernel RBF, dan 74.89% untuk SVM kernel sigmoid. Sehingga dapat menyimpulkan dari skenario ini, penggunaan SVM linear lebih efisien dibanding model SVM yang lain yang diuji dalam penelitian ini. Hal ini disebabkan metode SVM linear dapat dipisahkan secara linear pada ruang input, karena margin dapat menemukan pemisah dalam hyperplane sehingga dapat mengklasifikasi dengan lebih baik.

4.4 Analisa Hasil Penelitian

Dengan dilakukan beberapa skenario pengujian terhadap penelitian, dapat dilihat bahwa beberapa faktor yang telah dilakukan dalam alur sistem penelitian ini dapat mempengaruhi nilai performansi akurasi dan *f1-score*. Seperti pengaplikasian proses *preprocessing lemmatization* pada dataset dapat mempengaruhi nilai *f1-score* sebesar 1.08%. Juga pengaruh penggunaan dimensi yang lebih besar pada proses word2vec skip-gram *negative sampling*, yang dimana antara 300 dimensi dan 100 dimensi memiliki selisih nilai *f1-score* sebesar 4.12%. Serta pengaruh penggunaan model SVM, yang dimana nilai *f1-score* tertinggi berada dalam model SVM linear dengan nilai 78.74%. Hal tersebut membuktikan bahwa faktor-faktor tersebut dapat mempengaruhi hasil akhir penelitian ini.

5. Kesimpulan

Berdasarkan hasil penelitian yang dilakukan, kesimpulan yang dapat diambil dari penelitian ini sebagai berikut. Metode ekstraksi fitur word2vec yang digabungkan dengan klasifikasi SVM dengan jumlah data sebanyak 10.000 ulasan. Dapat menyimpulkan tahap *preprocessing*, dimensi word2vec, dan model klasifikasi mempengaruhi hasil performansi pada analisis sentimen level aspek pada ulasan film. Hasil di setiap skenario terbaik yaitu, menggunakan tahap *preprocessing lemmatization*, word2vec menggunakan 300 dimensi, dan model klasifikasi SVM linier dengan hasil terbaik yaitu sebesar 78,74% *f1-score*. Penggunaan jumlah dimensi word2vec yang ada pada data train mempengaruhi akurasi dari klasifikasi. Hal ini disebabkan karena semakin tinggi dimensi, maka variasi dari *vocabulary* yang dibentuk oleh model word2vec akan semakin tinggi variasinya, sehingga akurasi meningkat.

Saran yang dapat dilakukan penelitian selanjutnya adalah meningkatkan penggunaan dimensi pada word2vec yang dimana berkemungkinan meningkatkan akurasi, dan juga lebih berfokus pada tahap *preprocessing*, yaitu memperbaiki kamus kata dan melakukan pengecekan terhadap list *stopword* yang digunakan untuk menghindari kata yang mempengaruhi performansi. Serta membangun sistem yang lebih efisien demi mengurangi *running time* program.

REFERENSI

- [1] Rizky, M. Y., & Stellarosa, Y. (n.d.). Prefensi Penonton Terhadap Film Indonesia. *Journal of communication Studies*.
- [2] G., V., & RM., C. (2012). Sentiment Analysis and Opinion Mining: A Survey. *International Journal of Advanced Research in Computer Science and Software Engineering*.
- [3] Widyaningtyas, C.W., Adiwijaya, A. dan Al Faraby, S., (2018). Klasifikasi *Sentiment Analysis* pada Review Film Berbahasa Inggris dengan Menggunakan Metode Doc2Vec dan Support Vector Machine (SVM).
- [4] Acosta, J., Lamaute, N., Luo, M., Finkelstein, E., & Andreea, C. (2017). *Sentiment analysis* of Twitter Messages Using Word2Vec. *CSIS*, 7.
- [5] Liu, B. (2012). Sentiment Analysis and Opinion Mining Synthesis *Lectures on Human Language Technologies*.
- [6] Nawangsari, R.P., Kusumaningrum, R. and Wibowo, A.(2019). Word2Vec for Indonesian Sentiment Analysis towards Hotel Reviews: An Evaluation Study
- [7] Rahmawati, F.F dan Sibaroni., (2019).Multi-aspek pada Destinasi Pariwisata Yogyakarta Menggunakan Support Vector Machine dan Particle Swarm Optimization sebagai Seleksi Fitur.
- [8] Fauzi,M.A., (2019). *Word2Vec model for sentiment analysis of product review* in Indonesia language. *International Journal of Electrical and Computer Engineering*, 9(1), p.525.
- [9] Vijayarani, S., Ilamathi, M.J. and Nithya, M., (2015). Preprocessing techniques for text mining-an overview. *International Journal of Computer Science & Communication Networks*, 5(1), pp.7-16.
- [10] Mikolov, T., Chen, K., Corrado, G. dan Dean, J., (2013). Efficient estimation of word representations in vector space. arXiv preprint arXiv:1301.3781.
- [11] Mikolov, T., Sutskever, I., Chen, K., Corrado, G. dan Dean, J., (2013). Distributed representations of words and phrases and their compositionality. arXiv preprint arXiv:1310.4546.
- [12] Purbalaksono, M.D., (2019). Skip-Gram Negative Sample for Word Embedding in Indonesian-Translation Text Classification (Doctoral dissertation, Telkom University).
- [13] Honakan, H., Adiwijaya, A. and Al Faraby, S., (2018). Analisis Dan Implementasi Support Vector Machine Dengan String Kernel Dalam Melakukan Klasifikasi Berita Berbahasa Indonesia. *Proceedings of Engineering*, 5(1).
- [14] Putri, L.R., Mubarak, M.S., Adiwijaya(2017) Klasifikasi Sentimen Ulasan Buku Berbahasa Inggris Menggunakan *Information Gain* dan *Naive Bayes*
- [15] Juliansyah, F.R., Adiwijaya., Purbalaksono, M.D., (2020) Analisis Sentimen Menggunakan Metode KNN dengan *Negation Handling* pada Data *Movie Review*