

Deteksi *Fake Reviews* Menggunakan *Support Vector Machine*

Bety Elysabeth Pasaribu¹, Anisa Herdiani², Widi Astuti³

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

¹betyelysabeth@students.telkomuniversity.ac.id, ²anisaherdiani@telkomuniversity.ac.id,

³widiwdu@telkomuniversity.ac.id

Abstrak

Maraknya berbagai *e-commerce* menjadikan calon pembeli semakin selektif sehingga bergantung pada *review* yang ditinggalkan oleh pembeli sebelumnya untuk menentukan keputusan membeli suatu produk. Banyaknya *review*, baik itu yang bersifat positif atau negatif, sangat mempengaruhi sisi mana yang dapat dipercaya. Jika *review* yang dibaca tidak nyata atau disebut *fake review* maka akan merugikan baik sisi penjual ataupun sisi pembeli. Untuk itu, perlu dilakukan analisis untuk mendeteksi *fake review* pada kumpulan *review* produk. Penelitian ini dilakukan dengan pendekatan lima kelas *feature* yaitu *sentiment feature*, *personal feature*, *brand-only feature*, *content feature*, dan *metadata feature* dengan menggunakan metode klasifikasi *Support Vector Machine*. Pada penelitian ini dibandingkan antara *SentiwordNet* dan *SenticNet* untuk mendapatkan ekstraksi *sentiment* mana yang lebih baik. Pada penelitian ini juga dilakukan pemilihan dan penggabungan *feature*, serta *tuning* parameter dan jenis *kernel* pada SVM apakah akan memengaruhi sistem. Hasil terbaik diperoleh akurasi sebesar 74,46%. Dari hasil penelitian ini diperoleh bahwa *SenticNet* lebih baik daripada *SentiwordNet*, kemudian *tuning* parameter serta pemilihan jenis *kernel* pada SVM bisa mendapatkan hasil yang optimal, serta penggunaan *sentiment feature* sangat mempengaruhi sistem untuk deteksi *fake review*.

Kata kunci: *fake reviews*, *support vector machine*, *feature*, *sentiwordnet*, *senticnet*

Abstract

Lot of various e-commerce makes prospective buyers more selective so that it relies on reviews left by The rise of various e-commerce makes prospective buyers more selective so that it relies on reviews left by previous buyers to determine the decision to buy a product. The number of reviews, both positive and negative, greatly influences which side can be trusted. If the review that is read is not real or is called a fake review, it will harm both of the seller and the buyer side. For this reason, an analysis is needed to detect fake reviews on a collection of product reviews. This research was approached with a five-class features named sentiment features, personal features, brand-only features, content feature, and metadata feature using the Support Vector Machine classification method. This research compares between SentiwordNet and SenticNet to get which sentiment extraction is better. This research also carried out to analyze whether the differences in the use of SentiwordNet and SenticNet, the selection and integration of features, and changes in parameters also choosing kernel in SVM will affect the system. The best results obtained an accuracy of 74,46%. From the results of this study, it was found that SenticNet is better than SentiwordNet, then tuning SVM parameters can get optimal results, also using sentiment feature affect the system for detecting fake review.

Keywords: *fake reviews*, *support vector machine*, *features*, *sentiwordnet*, *senticnet*

1. Pendahuluan

Latar Belakang

Saat ini, semakin banyak aplikasi maupun website untuk kegiatan *e-commerce* yang dapat dipilih konsumen untuk mendapatkan produk yang diinginkan. Namun, sebelum memutuskan untuk membeli suatu barang atau produk, calon pembeli sering melihat dahulu *review* atau ulasan yang ditinggalkan oleh pembeli sebelumnya untuk dapat melihat apakah produk yang akan dibeli berkualitas atau tidak berdasarkan pengalaman orang lain yang menuliskan *review* produk tersebut, bahkan untuk satu produk yang sama, calon pembeli perlu melihat berbagai *review* di *platform* yang berbeda. Akan tetapi, *review* yang dibagikan juga belum tentu sepenuhnya benar dan dapat dipercaya atau dapat dikatakan juga sebagai *review* palsu. Bahkan,

adanya organisasi ataupun kelompok yang dibayar untuk memberikan *review* yang bersifat menjatuhkan atau menaikkan pamor produk [1].

Salah satu cara untuk mengidentifikasi *fake review* ialah apabila keseluruhan *review* berisikan kalimat positif atau kalimat negatif maka dapat dicurigai sebagai *review* palsu [2]. Apabila semua kalimat bersifat positif yang diberikan untuk suatu produk, maka dapat dianggap sebagai pemasaran produk untuk mengecoh pembaca dengan mempromosikan suatu produk untuk menaikkan reputasi produk. Sedangkan, jika semua kalimat negatif yang diberikan untuk suatu produk maka dapat dianggap sebagai ulasan yang bersifat menjatuhkan produk tersebut untuk merusak reputasinya. Hal ini tentu sangat merugikan para pembeli karena termakan oleh *review* palsu yang beredar dan mendapatkan produk yang tidak bagus. Dari sisi penjual pun akan sangat dirugikan apabila banyak mendapat *review* palsu karena akan menjatuhkan usaha penjual. Karena itu sudah banyak peneliti yang menganalisis *review* palsu atau disebut juga sebagai *opinion spammer* untuk mempelajari karakteristik *review* palsu tersebut.

Penelitian untuk *fake review* telah dilakukan sebelumnya oleh [3], namun pelabelan yang digunakan adalah teknik *duplicate and near duplicate* yang dinilai tidak tepat sasaran karena *review* yang duplikat belum tentu merupakan *spam* atau sebaliknya [2]. Selanjutnya, penelitian dengan menggunakan Metode Naïve Bayes berbasis *feature* yang menghasilkan akurasi sebesar 83,33%, penelitian ini menggunakan NLP dan sumber *lexicon* untuk memproses *feature* serta menggabungkan semua fitur untuk mendapatkan akurasi yang baik [4], namun pada penelitian ini kelas *brand-only feature* tidak disertakan. Pada Tugas Akhir ini, dilakukan penelitian untuk mendeteksi *fake review* dengan menggunakan POS Tagging, SentiWordNet, SenticNet, dan metode *lookup* untuk mengekstraksi fitur dan proses klasifikasi menggunakan *Support Vector Machine* (SVM). Dataset yang digunakan telah dikumpulkan dari Amazon.com [3]. Selanjutnya dilakukan pelabelan secara manual seperti yang dilakukan oleh [2,4]. Setelah itu, dilakukan *preprocessing* yang terdiri dari *case folding*, *removing stopword*, *cleansing*, *POS Tagging*. Setelah melakukan *preprocessing*, dilanjutkan dengan menghitung jumlah setiap kelas fitur dengan menggunakan SentiWordNet, SenticNet, dan metode *lookup*. Selanjutnya, setiap nilai yang didapatkan menjadi masukan pada proses klasifikasi menggunakan metode SVM. Metode ini digunakan karena dinilai memiliki performansi yang baik dan dapat menghasilkan akurasi yang tinggi [1,5]. Metode ini juga membantu mengidentifikasi *feature* apa yang lebih mempengaruhi dalam pengklasifikasian untuk mendeteksi *fake review*.

Topik dan Batasannya

Rumusan masalah pada penelitian tugas akhir ini adalah seputar deteksi *fake review* dengan menggunakan *sentiment feature*, *personal feature*, *brand-only feature*, *content feature* serta *metadata feature* untuk mengetahui bagaimana hasil performansi sistem yang dibangun dan *feature* apa yang paling mempengaruhi sistem.

Pada penelitian tugas akhir ini terdapat beberapa batasan yakni dataset berbahasa Inggris yang pelabelannya dilakukan secara manual dengan jumlah *review* sebesar 940 *review* dan tidak dapat digunakan untuk level *chunk*. Batasan untuk setiap *review* yang digunakan dalam penelitian ini maksimal untuk 1500 karakter.

Tujuan

Tujuan dari penelitian tugas akhir ini adalah untuk mengetahui bagaimana hasil performansi sistem yang dibangun dan kelas *feature* apa yang paling mempengaruhi sistem dengan menggunakan metode klasifikasi *Support Vector Machine*.

Organisasi Tulisan

Organisasi penulisan tugas akhir ini dimulai dari Pendahuluan dimana dijelaskan mengenai topik permasalahan tugas akhir yang diambil dan apa yang dilakukan secara garis besar terkait deteksi *fake review*. Dari Pendahuluan dilanjutkan dengan bagian Studi Terkait untuk menjelaskan mengenai studi literatur terkait topik permasalahan tugas akhir. Selanjutnya, pada bagian Sistem yang Dibangun, menjelaskan setiap proses yang dilakukan untuk membangun sistem, dimulai dari penjelasan tentang dataset, *praproses*, ekstraksi fitur sampai klasifikasi SVM. Selanjutnya, Hasil Pengujian dan Analisis untuk menjelaskan serta memberikan analisis dari hasil yang didapatkan dari sistem yang dibangun. Setelah itu, diberikan kesimpulan dari analisis yang diberikan serta saran terkait untuk keperluan penelitian selanjutnya. Bagian terakhir yaitu Daftar Pustaka sebagai susunan daftar dari setiap referensi yang digunakan dalam penyelesaian tugas akhir. Kemudian Lampiran untuk menyertakan penjelasan yang lebih detil mengenai data ataupun hasil uji coba.

2. Studi Terkait

2.1 Opinion Mining

Opinion mining dan *sentiment analysis* memiliki wilayah studi yang sama, sehingga *opinion mining* juga bisa disebut sebagai *sentiment analysis*, begitu juga sebaliknya [10]. Sentimen analisis yang juga disebut

sebagai *opinion mining* adalah bidang studi yang menganalisis opini, sentimen, evaluasi, pujian, sikap, dan emosi terhadap entitas seperti produk, jasa, organisasi, individual, isu, acara, topik dan atribut-atributnya [11]. Berdasarkan [12], ada beberapa topik yang berfokus pada *sentiment analysis* atau *opinion mining*, salah satunya ialah:

2.1.1 Opinion Spam

Opinion spam merujuk kepada opini palsu yang mencoba mengecoh pembaca atau sistem otomatis dengan memberikan opini positif yang tidak layak kepada target objek untuk mempromosikan objek tersebut atau dengan memberi opini negatif jahat untuk merusak reputasinya. Mendeteksi spam seperti ini sangat penting untuk aplikasi. Salah satu tipe *opinion spam* ialah *untruthful opinion* [3], yang sengaja menyesatkan pembaca dengan memberikan ulasan positif yang tidak pantas ke objek yang ditargetkan untuk mempromosikan objek atau memberikan ulasan negatif yang tidak adil untuk objek lain untuk merusak reputasinya. *Untruthful reviews* secara umum juga disebut sebagai *fake reviews* atau *bogus review*. Untuk mendeteksi *untruthful review* digunakan fitur berbasis *review centric* yaitu yang berhubungan dengan *review* untuk mengenali *spam/fake review* [3]. Terdapat 5 jenis fitur utama yang digunakan, yaitu *sentiment feature*, *personal feature*, *brand-only feature*, *content feature*, dan *metadata feature* berdasarkan peneliti [2]:

a. Sentiment feature (SF)

- *Positive vs Negative*

Jika *review* hanya mengekspresikan *sentiment* positif atau *sentiment* negatif pada produk, maka cenderung *spam/fake*. Karena *review* yang benar akan mengekspresikan kedua *sentiment* [2].

b. Personal feature (PF)

- *First Person vs Second Person*

Ditemukan bahwa dalam *review* palsu, sering ditemukan ““you” should do something” (kamu harus bla bla), daripada bagaimana ““I” experience” (pengalaman “saya”) [2]. Karena itu, *review* yang jujur akan lebih banyak menggunakan *1st pronoun* dibandingkan dengan *2nd pronoun*.

c. Brand-only feature (PRF)

Dalam suatu *review*, seseorang mungkin untuk memberikan *review* sebagai contoh “*HP is the best brand*” atau “*Canon is better than Fuji*”. Hal ini tidak mengekspresikan opini seseorang pada suatu produk melainkan pada suatu *brand*. Ini juga menunjukkan adanya *bias* pada *review* atau opini tersebut [2,3]

d. Content Feature (CF)

Content feature yang digunakan dalam penelitian ini adalah *helpful feedback*, *feedback*, *length title*, *length review* [2].

e. Metadata Feature (MF)

Metadata feature yang digunakan dalam penelitian ini adalah *rating* yang diberikan [2].

2.2 POS Tagging

POS Tagging adalah identifikasi semua jenis kata dalam konteks. POS Tagging adalah operasi yang penting dan dibutuhkan oleh semua sistem NLP (*Natural Language Processing*). Langkah ini merupakan awal dari analisis banyak *parsing* [16]. Tentang pengaplikasiannya, Berikut ini contoh dari POS Tagging:

My dog also like eating sausage

Diproses dengan POS Tagging akan memiliki hasil berikut:

My/PRP\$ dog/NN also/RB likes/VBZ eating/VBG sausage/NN

2.3 SentiWordNet

SentiWordNet adalah sumber leksikal informasi sentimen untuk *opinion mining*. *WordNet* disebut juga sebagai database leksikal untuk Bahasa Inggris. Pada *WordNet*, kata dikelompokkan ke dalam set sinonim yang disebut *synset*. *SentiWordNet* merupakan hasil anotasi otomatis dari semua *synset* pada *WordNet* diberikan skor numerik untuk tiga sentimen yaitu positif, negative dan objektif untuk menjelaskan seberapa positif, negatif dan objektif suatu term yang ada pada *synset*. Ketiganya diberi range skor dari 0.0 sampai 1.0 dan jumlahnya adalah 1.0 untuk setiap *synset* [13].

2.4 SenticNet

SenticNet merupakan pengukuran semantik menggunakan pengetahuan dengan memanfaatkan kedua teknik web semantik dengan *artificial intelligence* untuk mengenali, memroses, dan menginterpretasikan opini NLP dari seluruh Web. Untuk mengidentifikasi, *SenticNet* memberikan skor antara -1 sampai 1 [16] dimana nilai dibawah 0 dianggap sebagai kata negatif, sedangkan kata dengan skor diatas 0 dianggap kata positif.

2.5 Support Vector Machine

SVM dalam pembelajaran mesin adalah model pembelajaran *supervised* dengan mempelajari algoritma yang terkait untuk menguji data dan mengidentifikasi pola, yang mana digunakan untuk regresi dan analisis klasifikasi [5]. Dua kelas data yang digambarkan sebagai lingkaran dan padat titik-titik yang disajikan angka

ini. Secara intuitif diamati, ada banyak keputusan *hyperplanes* yang dapat digunakan untuk memisahkan kedua kelompok data [14]. Upaya mencari *hyperplane* optimal ini merupakan inti dari proses pada SVM [15].

Data yang tersedia dinotasikan sebagai $\vec{x}_i \in \mathbb{R}^d$ sedangkan label masing-masing dinotasikan $y_i \in \{-1, +1\}$ untuk $i=1,2,\dots,l$, yang mana l adalah banyaknya data. Diasumsikan kedua kelas -1 dan +1 dapat terpisah secara sempurna oleh *hyperplane* berdimensi d , yang didefinisikan

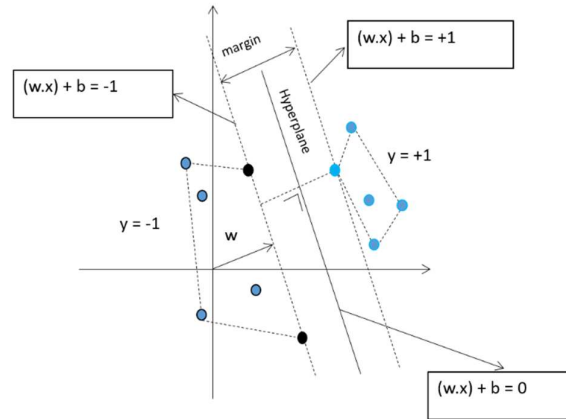
$$\vec{w} \cdot \vec{x} + b = 0 \quad (1)$$

Pattern $\vec{x}_i$ yang termasuk kelas -1 (sampel negatif) dapat dirumuskan sebagai pattern yang memenuhi pertidaksamaan

$$\vec{w} \cdot \vec{x} + b \leq -1 \quad (2)$$

Sedangkan pattern $\vec{x}_i$ yang termasuk kelas +1 (sampel positif)

$$\vec{w} \cdot \vec{x} + b \geq +1 \quad (3)$$



Gambar 3.1 Kernel Linear [21]

Margin terbesar dapat ditemukan dengan memaksimalkan nilai jarak antara *hyperplane* dan titik terdekatnya, yaitu $1/\|\vec{w}\|$. Hal ini dapat dirumuskan sebagai Quadratic Programming (QP) problem, yaitu mencari titik minimal persamaan (4), dengan memperhatikan constraint persamaan (5).

$$\min_{\vec{w}} \tau(\vec{w}) = \frac{1}{2} \|\vec{w}\|^2 \quad (4)$$

$$y_i(\vec{x}_i \cdot \vec{w}) - 1 \geq 0, \forall i \quad (5)$$

Problem ini dapat dipecahkan dengan berbagai teknik komputasi, di antaranya Lagrange Multiplier.

$$L(\vec{w}, b, \alpha) = \frac{1}{2} \|\vec{w}\|^2 - \sum_{i=1}^l \alpha_i (y_i(\vec{x}_i \cdot \vec{w} + b) - 1) \quad (6)$$

$(i = 1, 2, \dots, l)$

α_i adalah Lagrange multipliers, yang bernilai nol atau positif ($\alpha_i \geq 0$). Nilai optimal dari persamaan (6) dapat dihitung dengan meminimalkan L terhadap $\vec{w}$ dan b , dan memaksimal L terhadap α_i . Dengan memperhatikan sifat bahwa pada titik optimal gradient $L=0$, persamaan (6) dapat dimodifikasi sebagai maksimalisasi problem yang hanya mengandung α_i saja, sebagaimana persamaan (7) dibawah.

Maximize:

$$\sum_{i=1}^l \alpha_i - \frac{1}{2} \sum_{i,j=1}^l \alpha_i \alpha_j y_i y_j \vec{x}_i \cdot \vec{x}_j \quad (7)$$

Subject to;

$$\alpha_i \geq 0 \quad (i = 1, 2, \dots, l) \quad \sum_{i=1}^l \alpha_i y_i = 0 \quad (8)$$

Dari hasil perhitungan ini akan mengeluarkan α_i yang bernilai positif sehingga data yang berhubungan dengan α_i positif inilah yang disebut dengan *support vector*. Setelah menemukan *support vector*, dapat diketahui pula *hyperplane* optimum dengan mencari persamaan (9) dan (10)[15]:

$$\vec{w} = \sum_{i=1}^n \alpha_i y_i \vec{x}_i \quad (9)$$

$$b = y_i - \vec{w}^T \vec{x} \quad (10)$$

2.5.1 Soft Margin

Penjelasan di atas berlaku apabila berdasarkan asumsi bahwa kedua belah kelas dapat terpisah secara sempurna oleh *hyperplane*. Akan tetap, pada umumnya kedua buah kelas tersebut tidak dapat terpisah secara

sempurna. Hal ini menyebabkan proses optimisasi tidak dapat diselesaikan, karena tidak ada w dan b yang memenuhi pertidaksamaan 5.

Untuk itu, pertidaksamaan 5 dimodifikasi dengan memasukkan slack variable ξ_i ($\xi_i \geq 0$), menjadi

$$y_i(w \cdot x_i + b) \geq 1 - \xi_i, \forall i \quad (11)$$

Minimize

$$\frac{1}{2} \|\vec{w}\|^2 + C \sum_{i=1}^N \xi_i \quad (12)$$

Parameter C bertugas mengontrol *tradeoff* antara *margin* dan *classification error*. Semakin besar nilai C , semakin besar *penalty* yang dikenakan untuk *tap classification error*.

2.6 Evaluasi Performansi

Evaluasi dilakukan dengan tujuan untuk mengetahui akurasi hasil klasifikasi data menggunakan metode SVM terhadap data. Hasil evaluasi merupakan kelas prediksi yang dibandingkan dengan kelas actual atau yang sebenarnya dengan *confussion matrix* pada Tabel 2.1:

Tabel 2.1 Confusion matrix[20]

CLASS	Actual Positive	Actual Negative
Predicted Positive	TP	FP
Predicted Negative	FN	TN

TP (*True Positive*) adalah kelas yang diprediksi oleh sistem positif dan sesuai dengan data aslinya, yaitu positif. TN (*True Negative*) adalah kelas yang diprediksi oleh sistem sebagai data negative dan sesuai dengan data aslinya yaitu negatif. FP (*False Positive*) adalah kelas yang diprediksi oleh sistem positif, tetapi data aslinya memiliki nilai negatif. FN (*False Negative*) adalah kelas yang diprediksi sistem negatif, tetapi data aslinya memiliki nilai positif [18]. Perhitungan evaluasi performansi untuk menguji hasil klasifikasi dari sistem yang telah dibangun adalah dengan melakukan penghitungan *accuracy*, *precision*, *recall* dan *F-measure*, dapat dilihat sebagai berikut:

$$Akurasi = \frac{TP+TN}{TP+FP+FN+TN} \quad (9) \quad Precision = \frac{TP}{TP+FP} \quad (10)$$

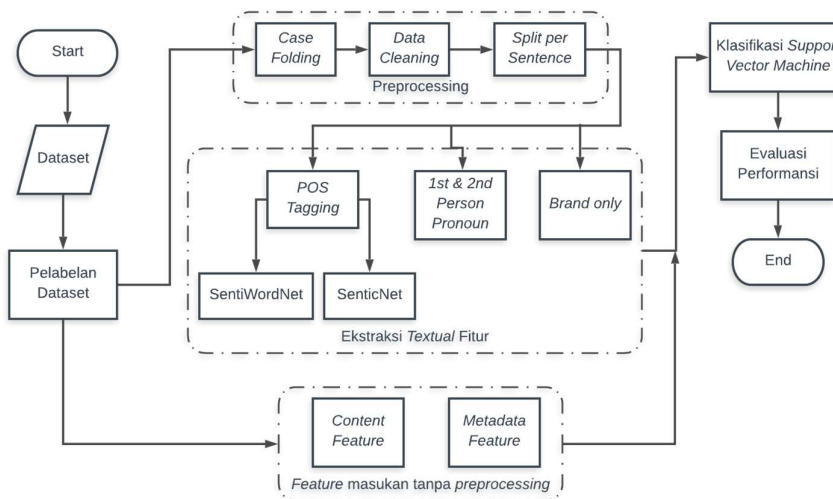
$$Recall = \frac{TP}{TP+FN} \quad (11) \quad F-Measure = 2 \times \frac{(Precision \times Recall)}{(Precision+Recall)} \quad (12)$$

Accuracy dapat menunjukkan tingkat keberhasilan sistem dalam melakukan klasifikasi. *Precision* adalah ukuran banyaknya dokumen yang ditemukan relevan, yaitu tingkat ketepatan antara hasil klasifikasi yang dilakukan oleh pengguna dengan hasil klasifikasi yang diberikan oleh sistem. *Recall* adalah tingkat keberhasilan sistem dalam menemukan kembali sebuah informasi. *F-Measure* adalah hasil rata-rata dari *precision* dan *recall*.

3. Sistem yang Dibangun

3.1 Gambaran Umum Sistem

Sistem pada tugas akhir ini dibangun menurut gambaran umum sistem pada Gambar 3.1 sebagai berikut:



Gambar 3.2 Gambaran umum sistem

Keterangan:



3.2 Pelabelan Dataset

Dataset yang digunakan dalam penelitian ini menggunakan data review konsumen dari Amazon yang dikumpulkan oleh Bing Liu [3]. Dataset yang digunakan berjumlah 940 data review dengan maksimal 1500 karakter. Pelabelan data review tersebut *fake* atau *nonfake* dilakukan secara manual oleh manusia dan dilakukan oleh 13 orang secara independen (data dibagi menjadi empat bagian, setiap bagian dilakukan pelabelan oleh minimal tiga orang) dipandu dengan beberapa instruksi¹ yang disarankan [2] dan digunakan juga dalam penelitian [4].

3.3 Preprocessing

Tahapan *preprocessing* ini merupakan proses awal yang berguna untuk menghilangkan *noise* secukupnya dan mempersiapkan teks menjadi data yang sesuai kebutuhan untuk diolah lebih lanjut. Tahapan *preprocessing* yang dilakukan ialah sebagai berikut:

1. *Case Folding*
Pada tahap ini akan diseragamkan semua huruf, yaitu mengubah semua huruf besar (*uppercase*) yang ada pada teks menjadi huruf kecil (*lowercase*).
2. *Data Cleaning*
Dalam tahap ini, akan dilakukan pembersihan untuk menghilangkan tanda baca dan yang selain huruf, seperti “”, @!#\$-%&*” dan simbol lainnya kecuali tanda titik (.).
3. *Split per Sentence*
Pada tahap ini akan dipisahkan setiap kalimat *review* dengan simbol titik “.”. Tahap ini diperlukan agar bisa diberikan penandaan yang sesuai saat proses POS Tagging. Sebagai contoh, kata “*mine*” akan berbeda penandaannya di kalimat “*This book is mine*” karena akan ditandai sebagai *posessive pronoun*, sedangkan pada kalimat “*Our last working mine ran dry three years ago*” akan ditandai sebagai *noun verb*.

Contoh tahapan *preprocessing* dapat dilihat pada Tabel 3.1 dibawah ini.

Tabel 3.1 Contoh Hasil *Case Folding*, *Data Cleaning*, dan *Split per Sentence*

Kalimat	<i>Within the last day or two, my LITEON 5005 can no longer recognize any rewriteable DVD+RW disk I put in it. It attempts to read it, but always comes back with either a 'FAIL' or an INVALID DISK message. As i only use this player/recorder in the living room (I have a comcast DVR on the main TV), I have not burned on this recorder very often. From the look of other posts, contacting technical support will do no good.</i>
Hasil setelah <i>Case Folding</i>, <i>Data Cleaning</i>, dan <i>Split per Sentence</i>	<i>within the last day or two my liteon can no longer recognize any rewriteable dvdrw disk i put in it.</i> <i>it attempts to read it but always comes back with either a fail or invalid disk message.</i> <i>as i only use this playerrecorder in the living room i have a comcast dvr on the main tv i have not burned on this recorder very often.</i> <i>from the look of other posts contacing technical support will do good.</i>

3.4 Ekstraksi Fitur

Sesudah kumpulan data melewati tahap *preprocessing*, akan dilakukan ekstraksi fitur pada konten *review* untuk bisa mempelajari *review* mana yang sudah dimanipulasi atau terdeteksi sebagai *fake review* [4,5]. Berikut proses ekstraksi fitur yang dilakukan dalam penelitian ini dengan berdasarkan *textual data centric review*[4]:

1. *POS Tagging*
Pada tahap ini akan diberikan penandaan pada setiap kata pada konten *review*. *Review* yang jujur cenderung memiliki banyak N(*Noun*), JJ(*Adjective*), IN(*Preposition/sub-conj*) dan

¹ <http://consumerist.com/2010/04/how-you-spot-fake-online-reviews.html>

DT(*Determiner*), sedangkan yang memanipulasi cenderung menghasilkan banyak V(*Verb*), RB(*Adverb*) dan PRP(*Personal Pronoun*) [7].

2. *SentiWordNet* dan *SenticNet*

Setelah dilakukan proses *POS Tagging*, selanjutnya akan diberikan skor pada kata yang terdapat dalam *SentiWordNet*. Apabila kata merupakan kata positif akan diberikan skor positif dan dijumlahkan setiap kata yang memiliki skor positif, demikian juga kata yang negatif akan diberikan skor negatif kemudian dijumlahkan untuk setiap kata yang memiliki skor negatif. *SenticNet* menjadi perbandingan untuk *SentiWordNet* agar dapat dilihat sumber leksikal informasi *sentiment* mana yang menghasilkan akurasi yang lebih baik. Pada penelitian ini, penggunaan *SenticNet* memiliki akurasi yang lebih baik.

3. *Personal Pronoun*

Pada proses ini, *personal pronoun* yang diikutsertakan hanya *1st* dan *2nd* *pronoun*. Setiap *review* yang mengandung *1st* *pronoun* (*i, me, myself*, dsb) akan dijumlahkan, begitu juga *2nd* *pronoun* (*you, your, yours*) akan dijumlahkan.

4. *Reviews on Brand only*

Pada proses ini, akan dijumlahkan berapa kali dalam satu *review* menyebutkan nama *brand* yang berbeda. Untuk dapat mengekstraksi fitur ini, akan digunakan *Dictionary lookup methods*, metode ini merupakan cara yang sederhana dan efektif untuk melakukan *string matching* [19]. *Dictionary lookup methods* mengambil sebuah *string* sebagai masukan dan mencoba mencari yang paling mirip di dalam kamus. Kamus daftar nama *brand* disesuaikan dengan kebutuhan penelitian ini.

Contoh tahapan ekstraksi fitur dapat dilihat pada Tabel 3.2 berikut.

Tabel 3.2 Contoh Hasil POS Tagging

Kalimat	<i>Within the last day or two, my LITEON 5005 can no longer recognize any rewriteable DVD+RW disk I put in it. It attempts to read it, but always comes back with either a 'FAIL' or an INVALID DISK message. As i only use this player/recorder in the living room (I have a comcast DVR on the main TV), I have not burned on this recorder very often. From the look of other posts, contacting technical support will do no good.</i>
Hasil POS Tagging + SenticNet	within[IN] (0,935), the[DT] (0,935), last[JJ] (0,067), day[NN] (0,761), or[CC] (0,000), two[CD] (-0,030), my[PRP\$] (0,000), liteon[NN] (0,000), can[MD] (0,000), no[RB] (0,000), longer[RB] (0,865), recognize[VB] (0,041), any[DT] (0,000), rewriteable[JJ] (0,000), dvdwr[NN] (0,000), disk[NN] (-0,840), i[FW] (0,000), put[VRB] (0,765), in[IN] (0,157), it[PRP] (0,000), it[PRP] (0,000), attempts[VBZ] (0,000), to[TO] (0,789), read[VB] (-0,020), it[PRP] (0,000), but[CC] (0,000), always[RB] (0,000), comes[VBZ] (0,000), back[RP] (0,823), with[IN] (0,000), either[CC] (0,000), a[DT] (0,000), fail[VRB] (-0,870), or[CC] (0,000), an[DT] (0,000), invalid[JJ] (-0,790), disk[NN] (-0,840), message[NN] (0,162),
Hasil POS Tagging + SentiWordNet	within[IN] (0,000), the[DT] (0,000), last[JJ] (0,000), day[NN] (0,500), or[CC] (0,000), two[CD] (0,000), my[PRP\$] (0,000), liteon[NN] (0,000), can[MD] (0,000), no[RB] (-0,750), longer[RB] (-0,000), recognize[VB] (1,500), any[DT] (0,000), rewriteable[JJ] (0,000), dvdwr[NN] (0,000), disk[NN] (-0,000), i[FW] (0,000), put[VRB] (0,125), in[IN] (0,000), it[PRP] (0,000), it[PRP] (0,000), attempts[VBZ] (0,000), to[TO] (0,000), read[VB] (0,125), it[PRP] (0,000), but[CC] (0,000), always[RB] (-0,250), comes[VBZ] (0,000), back[RP] (-0,000), with[IN] (0,000), either[CC] (0,000), a[DT] (0,000), fail[VRB] (-3,125), or[CC] (0,000), an[DT] (0,000), invalid[JJ] (0,000), disk[NN] (-0,000), message[NN] (-0,000),

Pada Tabel 3.2, hasil pelabelan dengan SentiWordNet dan SenticNet dapat diperhatikan bahwa, SenticNet lebih banyak memberikan skor pada setiap kata. Contohnya untuk kata “*invalid*” dengan label JJ (*Adjective*) pada SentiWordNet tidak ada skor yang diberikan, sedangkan SenticNet memberikan skor sebesar -0,790 yang berarti kata ini masuk ke dalam kategori kata negatif dengan skor 0,790. Hal ini membuktikan bahwa penggunaan SenticNet lebih baik daripada SentiWordNet. Untuk selanjutnya,

Tabel 3.3 Contoh Hasil Ekstraksi Fitur

	<i>Sentiment Feature</i> (SF)	<i>Personal Feature</i> (PF)	<i>Brand-Only Feature</i> (PRF)
--	-------------------------------	------------------------------	---------------------------------

	Jumlah Kata Positif	Jumlah Kata Negatif	Jumlah 1 st Pronoun	Jumlah 2 nd Pronoun	Jumlah nama <i>brand</i> yang disebutkan
Hasil Ekstraksi	5,64	16,345	4,00	0,00	2,00

Tabel 3.4 Feature Masukan Tanpa Preprocessing

	Content Feature (CF)				Metadata Feature (MF)
	Helpful Feedback	Feedback	Length Title	Length Review	Rating
Hasil	6	6	24	418	1

3.5 Klasifikasi Menggunakan Support Vector Machine

Pada proses ini dilakukan klasifikasi menggunakan metode SVM untuk mendapatkan hasil akhir dari pembuatan sistem. Hasil ekstraksi fitur merupakan masukan yang digunakan dalam klasifikasi SVM sebagai metode *supervised* untuk mempelajari *fake review*.

Tahapan yang dilakukan pada pembangunan model klasifikasi SVM ini ialah;

1. Pembagian dataset dilakukan dengan menggunakan *10 Fold Cross Validation*. Sistem akan melakukan 10 kali *training* dan validasi, dimana setiap eksperimen menggunakan data partisi ke-1,2,3, dst sampai 10 bergantian disetiap batch *cross validation*-nya sebagai data uji dan memanfaatkan sisa partisi lainnya sebagai data latih [8].
2. Pada penelitian ini, dilakukan uji coba antara *kernel* Linear, Polinomial dan RBF untuk mengetahui jenis *kernel* mana yang mendapatkan akurasi terbaik.
3. Berdasarkan peneliti [9], untuk mendapatkan *gamma*, *C*, dan *epsilon* yang optimal dapat dilakukan pendekatan *grid-search* untuk *tuning* parameter antara 2^{-15} sampai 2^{15} . Namun, pada penelitian tugas akhir ini hanya dilakukan *tuning* parameter dari 2^{-5} sampai 2^5 .

4. Hasil Pengujian dan Analisis

4.1 Hasil Pengujian Setiap Kernel

Tabel 4.1 merupakan hasil komparasi setiap *kernel* yang diujikan pada data. Pengujian setiap *kernel* ini dilakukan untuk melihat perbedaan nilai akurasi yang dihasilkan dari penggunaan jenis *kernel*. Berdasarkan hasil yang diberikan, *kernel* Linear memiliki akurasi yang lebih tinggi sebesar 74,46 % dengan nilai $C=2^{-3}$. Sehingga dapat disimpulkan bahwa *kernel* yang lebih baik digunakan untuk penelitian deteksi *fake review* ini adalah *kernel* Linear.

Tabel 4.1 Komparasi Akurasi Setiap Kernel

Jenis Kernel	Akurasi
Linear	74,46%
Polinomial	65,75%
RBF	69,83%

4.2 Hasil Pengujian POS Tagging+SentiWordNet vs POS Tagging+SenticNet

Tabel 4.2 merupakan hasil komparasi penggunaan *POS Tagging* dengan penggunaan sumber leksikal informasi sentimen diajukan yaitu *SentiWordNet* dan *SenticNet*. Berdasarkan hasil yang diberikan, dapat diketahui bahwa penggunaan *POS Tagging* + *SenticNet* untuk *sentiment feature* dapat menghasilkan akurasi yang lebih tinggi dibandingkan dengan menggunakan *POS Tagging* + *SentiWordNet*. Hal ini disebabkan karena, *SenticNet* lebih mampu untuk memberikan skor pada kata yang tidak dapat diberikan oleh *SentiWordNet*.

Tabel 4.2 Komparasi Akurasi POS Tagging+SentiWordNet vs POS Tagging+SenticNet

Jenis Pengujian	Akurasi
<i>Sentiment Feature</i> + <i>POS Tagging</i> + <i>SentiWordNet</i>	67,19%
<i>Sentiment Feature</i> + <i>POS Tagging</i> + <i>SenticNet</i>	68,53%

4.3 Hasil Pengujian Setiap Fitur

Tabel 4.3 merupakan hasil pengujian untuk setiap *feature* dan kombinasi dengan *feature* yang lain. Kombinasi *Textual Feature* yang terdiri dari *sentiment feature*, *personal feature*, dan *brand-only feature* dengan *content feature* dan *metadata feature* mendapatkan hasil akurasi yang lebih baik sebesar 74,46%. Dapat diperhatikan bahwa pengujian hanya untuk kelas *sentiment feature* memberikan akurasi sebesar 73,63%. Hal ini menunjukkan bahwa kelas *sentiment feature* sangat berpengaruh dalam menentukan

suatu review merupakan *fake* atau bukan. Dari hasil pengujian, analisis yang dapat diberikan ialah semakin banyak jenis fitur yang digunakan, maka akan mampu mendapatkan hasil performansi yang lebih baik. Perlu diketahui juga bahwa nilai presisi dan recall pada penelitian ini rendah dipengaruhi oleh persebaran kelas data yang *imbalance*.

Tabel 4.3 Hasil Pengujian Setiap Fitur

Jenis Fitur	Akurasi%	Presisi%	Recall%	F-Measure%	Jenis Fitur	Akurasi%	Presisi%	Recall%	F-Measure%
SF	68,53	21,37	14,69	16,74	SF+PF+CF	73,69	72,8	46,04	52,8
PF	64,19	14,50	17,16	13,78	SF+PF+MF	71,28	60,19	41,81	45,63
PRF	65,50	0,00	0,00	0,00	SF+PRF+CF	71,45	71,6	29,46	39,9
CF	69,63	54,88	35,07	36,70	SF+PRF+MF	72,52	62,09	44,64	48,70
MF	65,75	0,00	0,00	0,00	SF+CF+MF	71,05	57,88	45,33	47,16
SF + PF	72,36	66,67	37,41	42,77	PF+PF+CF	71,31	74,26	35,68	43,91
SF + PRF	70,73	50,89	42,25	43,54	PF+PRF+MF	65,75	0,00	0,00	0,00
SF+CF	72,89	67,66	49,82	53,50	PF+CF+MF	72,49	68,06	33,76	41,99
SF+MF	68,73	42,11	33,52	34,15	PRF+CF+MF	73,60	70,40	48,90	55,35
PF + PRF	65,48	2,00	0,24	0,43	SF+PF+CF+MF	72,09	68,05	41,91	47,52
PF+CF	71,86	67,70	48,14	51,09	SF+PRF+CF+MF	74,09	77,05	34,47	44,88
PF+MF	65,53	31,90	21,50	23,59	PF+PRF+CF+MF	73,44	71,37	44,62	51,31
PRF+CF	74,30	73,27	42,10	52,00	ATF+CF	70,84	75,06	29,13	37,01
PRF+MF	65,66	0,00	0,00	0,00	ATF+MF	73,78	71,46	45,79	52,43
CF+MF	69,39	53,73	20,35	25,68	ATF+CF+MF	74,46	77,69	38,68	48,75
SF+PF+PRF(ATF)	73,63	65,46	46,96	51,33					

5. Kesimpulan dan Saran

5.1 Kesimpulan

Berdasarkan hasil pengujian dan analisis yang telah dilakukan, dapat disimpulkan:

1. Klasifikasi teks *review* menggunakan lima kelas feature dan dengan metode klasifikasi SVM dapat memberikan akurasi yang lebih baik ketika semua kelas feature yang diberikan.
2. Selain kelas *feature*, yang dapat mempengaruhi hasil performansi adalah jenis *kernel* dan parameter yang digunakan untuk SVM sebagai metode klasifikasi. Hal ini sesuai dengan hasil akurasi yang didapatkan yaitu sebesar 74,46% untuk jenis *kernel* Linear dengan nilai $C=2^{-3}$ sebagai parameter terbaik. Penggunaan antara *SenticNet* dan *SentiWordNet* juga dapat mempengaruhi hasil performansi. Hal ini dikarenakan ada beberapa kata yang terlewat dan tidak tersedia dalam sumber *lexicon SentiWordNet*, sedangkan pada *SenticNet* tersedia, sehingga pemberian skor lebih efektif.
3. Jenis *feature* yang paling berpengaruh adalah *sentiment feature*. Hal ini sesuai dengan hasil yang didapatkan berdasarkan hasil performansi klasifikasi dengan uji coba satu jenis *feature*.

5.2 Saran

Saran yang didapatkan dari penelitian tugas akhir ini untuk pengembangan sistem selanjutnya, yaitu;

1. Menambahkan kelas *feature* baru
2. Menambahkan fitur kata dengan level *chunk*. Contohnya : “*The book is worth your money.*”, bisa memiliki dua makna jika terpisah, yaitu antara “*book*” yang “*worth*” atau “*money*” yang “*worth*”. Jenis kata-kata seperti ini sering ditemui dalam *marketing speak*.

Daftar Pustaka

- [1] Algotar, K., & Bansal, A. (2018). Detecting truthful and useful consumer reviews for products using opinion mining. *CEUR Workshop Proceedings, 2111*, 63–72.
- [2] Li, F., Huang, M., Yang, Y., & Zhu, X. (2011). Learning to identify review spam. *IJCAI International Joint Conference on Artificial Intelligence*, 2488–2493.
- [3] Jindal, N., Liu, B., & Street, S. M. (2008). "Opinion and Spam Analysis." Proceedings of First ACM International Conference on Web Search and Data Mining (WSDM-2008), Feb 11-12, 2008, USA.
- [4] Anggraeni, A. C., Baizal, Z. K. A., & Setiawan, E. B. (2014). Analisis Deteksi Fake Review Pada User Online Review Berbasis Feature Dengan Menggunakan Metode Naïve Bayes. *Universitas Telkom*, 1–8.
- [5] Elmurngi, E., & Gherbi, A. (2017). Detecting Fake Reviews through Sentiment Analysis Using Machine Learning Techniques. *DATA ANALYTICS 2017 : The Sixth International Conference on Data Analytics Detecting*, (c), 65–72.
- [6] Chowdhary, N. S., & Pandit, A. A. (2018). Fake Review Detection using Classification. *International Journal of Computer Applications*, 180(50), 16–21.
- [7] Dave, K., Lawrence, S., & Pennock, D. M. (2003). Mining the peanut gallery: Opinion extraction and semantic classification of product reviews. *Proceedings of the 12th International Conference on World Wide Web, WWW 2003*, 519–528.
- [8] P. Refaailzadeh, L. Tang and H. Liu, "Cross-Validation," Arizona State University, 2008
- [9] Chih-Wei Chih-Wei Hsu, Chih-Chung Chang, and C.-J. L. (2008). A Practical Guide to Support Vector Classification. *BJU International*, 101(1), 1396–400.
- [10] Pang, B., & Lee, L. (2016). Opinion mining and sentiment analysis. *World Journal of Gastroenterology*, 22(45), 9898–9908.
- [11] Liu, B. (2012). Sentiment Analysis and Opinion Mining, (May).
- [12] Liu, B. (2010). Sentiment Analysis and Subjectivity. Handbook of Natural Language Processing, Second Edition, (Editors: N. Indurkha and F. J. Damerau), 2010, 1–38.
- [13] Hamouda, A., & Rohaim, M. (2011). Reviews classification using sentiwordnet lexicon. *World Congress on Computer Science and Information Technology*.
- [14] Arifin, Y. T. (2016). Komparasi Fitur Seleksi Pada Algoritma Support Vector Machine Untuk Analisis Sentimen Review. *Jurnal Informatika (JI) UBSI*, 3(September), 191–199.
- [15] Nugroho, A. S., Witarto, A. B., & Handoko, D. (2003). Support Vector Machine –Teori dan Aplikasinya dalam Bioinformatika1–.
- [16] Cambria, E., Speer, R., Havasi, C., & Hussain, A. (2010). SenticNet: A publicly available semantic resource for opinion mining. *AAAI Fall Symposium - Technical Report, FS-10-02*, 14–18.
- [17] Long, N. H., Nghia, P. H. T., & Vuong, N. M. (2014). Opinion Spam Recognition Method for Online Reviews Using Ontological Features. *Tap Chí Khoa Học*, 61, 44–59.
- [18] Vaitheeswaran, G. (2016). Combining Lexicon and Machine Learning Method to Enhance the Accuracy of Sentiment Analysis on Big Data. 7(1), 306–311.
- [19] Haldar, R., & Mukhopadhyay, D. (2011). Levenshtein Distance Technique in Dictionary Lookup Methods: An Improved Approach. (January 2011).
- [20] Visa, S., Ramsay, B., Ralescu, A., & Van Der Knaap, E. (2011). Confusion matrix-based feature selection. *CEUR Workshop Proceedings*, 710, 120–127.
- [21] Cholissodin, I., Setiawan B.D. (2013). Sentiment Analysis Dokumen E-Complaint Kampus Menggunakan Additive Selected Kernel SVM. *Seminar Nasional Teknologi Informasi dan Aplikasinya (SNATIA)*.