

Analisis sentimen ujaran kebencian pemilihan presiden 2019 menggunakan algoritma Naïve Bayes

Muftia Chalida¹ dan M. Didik R. Wahyudi^{*2}

- 1 Teknik Informatika UIN Sunan Kalijaga
Jl. Marsda Adisucipto Yogyakarta
muftiachalida97@gmail.com
- 2 Teknik Informatika UIN Sunan Kalijaga
Jl. Marsda Adisucipto Yogyakarta
m.didik@uin-suka.ac.id

Abstrak

Media sosial dapat memberikan gambaran secara umum opini yang terjadi didalam masyarakat, termasuk dalam pemilihan presiden 2019. Hal ini mengakibatkan bahwa data yang terkumpul dari media sosial sangat menarik untuk dianalisa guna mengetahui bagaimana suatu opini yang terjadi di masyarakat. Pengumpulan data harus memakai metode tertentu agar menghasilkan keakuratan opini yang terjadi di masyarakat. Penelitian ini mempergunakan teknik pengumpulan data dengan metode multistage random, berdasarkan data dari situs semiocast terhadap keaktifan postingan twitter di beberapa kota besar di Indonesia, yaitu Jakarta, Bandung, Semarang, Surabaya, Yogyakarta. Pengambilan data berdasarkan kata kunci pilpres 2019 yang dilakukan di beberapa kota di Indonesia diperoleh sebanyak 5055 data. Data ini kemudian di klasifikasi berdasarkan kategori ujaran kebencian dengan mempergunakan algoritma Naive Bayes Classifier dan pembobotan TF-IDF. Hasil yang diperoleh dari klasifikasi ini menunjukkan bahwa sentimen irrelevant sebanyak 11,3% dengan 573 data, sentimen negatif sebanyak 35,4% dengan 1786 data, sentimen netral sebanyak 26,7% sebanyak 1350 data dan sentimen positif sebanyak 26,6% sebanyak 1343 data. Sentimen negatif pada lima kota tersebut, memperoleh skor tertinggi dengan nilai sebesar 35,4%. Distribusi sentimen negatif pada lima kota yang dijadikan sampel menunjukkan bahwa

di Jakarta memiliki sentimen negatif sebesar 33,8%, Bandung sentimen negatif sebesar 65,4%. Surabaya sentimen positif sebesar 37,2 %, Yogyakarta dengan sentimen negatif sejumlah 51,8% dan Semarang dengan sentimen negatif 61,7%.

Kata Kunci analisis sentimen, naïve bayes classifier, multistage random, ujaran kebencian, pilpres 2019

1 Pendahuluan

Bercermin pada pelaksanaan pemilihan presiden tahun 2014, maraknya penyebaran isu yang berbau suku agama dan ras (SARA) dan ujaran kebencian diprediksi akan kembali terjadi pada pemilihan presiden tahun 2019. Penyebaran isu tersebut merupakan salah satu bentuk *black campaign* untuk menjatuhkan kredibilitas lawan politik di mata pemilih. Isu yang paling potensial untuk menyerang adalah sentimen agama. Penyebaran isu SARA dan ujaran kebencian merupakan salah satu strategi yang digunakan kelompok kepentingan tertentu untuk mencapai target yang diinginkan. Kegiatan ini ditujukan terutama untuk menurunkan

* Corresponding author.



kredibilitas guna mengurangi jumlah dukungan pihak lawan. Penyebaran isu dan ujaran kebencian tersebut, merupakan tindakan yang melanggar Undang-Undang Nomor 11 Tahun 2008 tentang Informasi dan Transaksi Elektronik serta Undang-Undang Nomor 40 Tahun 2008 tentang Penghapusan Diskriminasi Ras dan Etnis.

Sejalan dengan kemajuan zaman, penyebaran ujaran kebencian berupa *black campaign* di atas tidak hanya dapat dijumpai di dunia nyata, akan tetapi lebih terfokuskan di dunia maya, utamanya di media sosial, tidak terkecuali di twitter. Saat ini, media sosial tidak dapat dipungkiri memiliki pengaruh yang sangat besar dalam kehidupan. Di Indonesia, Twitter merupakan salah satu media sosial yang banyak digemari oleh masyarakat. Terlebih lagi, kemudahan yang disediakan oleh telepon seluler yang ada serta aplikasi yang mendukung. Hal ini membuat Indonesia menduduki peringkat ke enam sebagai negara dengan pengguna Twitter terbanyak, meski Amerika masih menjadi negara nomor satu untuk urusan Twitter. Twitter banyak digunakan untuk membahas mengenai kehidupan mereka, berbagi opini tentang berbagai topik dan membahas isu-isu yang terjadi pada saat ini. Format pesan yang bebas dan aksesibilitas dari berbagai platform yang mudah menjadikan pengguna internet cenderung untuk beralih dari blog atau milis ke layanan *microblogging* [1].

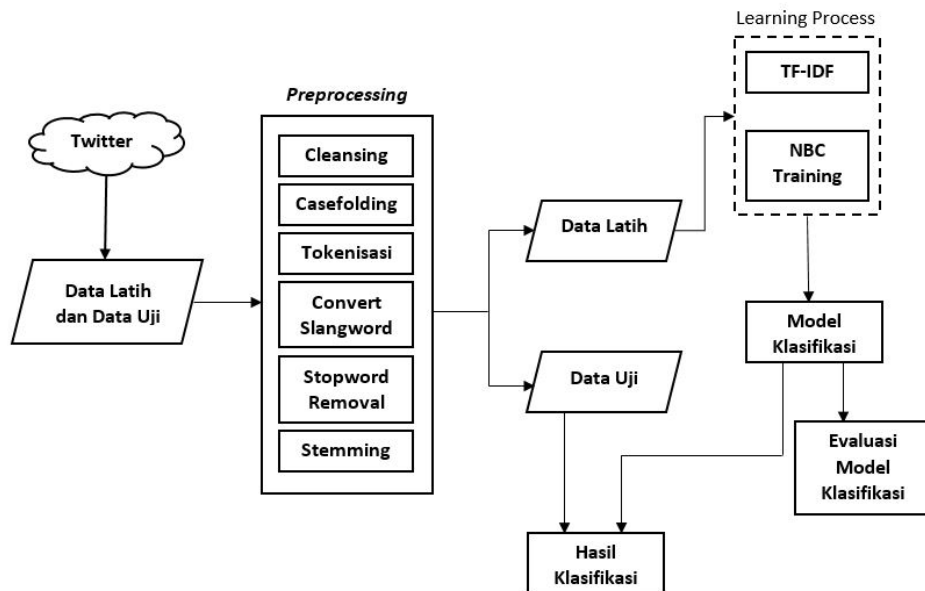
Hal tersebut menyebabkan semakin banyak pengguna yang melakukan posting tentang suatu produk dan layanan yang mereka gunakan, atau mengekspresikan pandangan mereka tentang politik maupun agama. Twitter sebagai salah satu situs *microblogging* dengan pengguna lebih dari 500 juta dan 400 juta tweet perhari [2]. Twitter dapat menjadi sumber data pendapat dan sentimen masyarakat dan tersebut dapat digunakan secara efisien untuk pemasaran atau studi sosial [3]. Oleh karena itu marak adanya *tweet* yang berhubungan dengan politik di tahun 2019 ini yang menjurus ke ujaran kebencian dapat digunakan untuk melakukan sosial media analitik (SMA). Analisis yang memuat mengenai kecenderungan informasi mengenai suatu topik apakah cenderung positif, negatif, netral ataukah irrelevant bisa dilakukan terhadap data tersebut Di Indonesia, Jakarta berdasarkan data dari SemioCast menempati urutan teratas dari 20 kota teraktif berdasarkan jumlah kicauan atau *tweet*. Sebanyak 1.058 miliar pada Juni lalu, kemudian diikuti dengan kota Bandung, Jawa Barat yang secara mengejutkan masuk dalam urutan keenam [4].

2 Metodologi

Penelitian ini akan melakukan analisis terhadap kecenderungan opini *tweet* mengenai ujaran kebencian di Pemilihan Presiden 2019. Metode pengumpulan data multistage random digunakan dalam penelitian ini. Data sampling didasarkan pada semioCast, dengan mengambil area meliputi Jakarta, Bandung, Semarang, Surabaya dan Yogyakarta. Data yang dianalisis adalah postingan *tweet* hasil dari pencarian dengan kata kunci pilpres2019. Data tersebut kemudian diklasifikasi menjadi empat kelas sentimen, yakni positif, negatif, netral, serta irrelevant. Metode yang akan digunakan untuk melakukan klasifikasi adalah Naïve Bayes Classifier (NBC).

NBC melakukan klasifikasi dan memprediksi probabilitas keanggotaan kelas suatu data *tuple* yang akan masuk ke kelas tertentu, sesuai dengan perhitungan probabilitas. NBC didasarkan pada teorema Bayes yang ditemukan oleh Thomas Bayes pada abad ke-18. Dalam suatu studi perbandingan, algoritma klasifikasi telah ditemukan sebagai bayesian sederhana atau yang biasa dikenal dengan NBC [5]. Fitur utama NBC adalah asumsi yang sangat kuat tentang independensi setiap kondisi atau peristiwa [6]. Algoritma NBC menggunakan metode pembelajaran mesin yang memanfaatkan probabilitas dan statistik yang disajikan oleh ilmuwan Inggris Thomas Bayes [7]. Perhitungan algoritma NBC dibandingkan

dengan algoritma klasifikasi lainnya lebih cepat karena hanya menguji probabilitas dengan menemukan kelas data pelatihan yang sama [8]. Selain itu, penggunaan metode Naive Bayes dianggap memiliki hasil akurasi yang cukup baik untuk klasifikasi *tweet* [9]. Tahapan penelitian yang dilakukan adalah diperlihatkan dalam Gambar 1.



■ **Gambar 1** Alur penelitian

2.1 Pengambilan data

Pengambilan data yang dipergunakan pada penelitian ini menggunakan metode *multistage random*, yaitu daerah sampling berupa 5 wilayah kota besar pulau jawa berdasarkan rujukan informasi kota dengan jumlah *tweet* terbanyak di Indonesia berdasarkan situs Semiocast, yaitu Jakarta, Bandung, Semarang, Surabaya, dan Yogyakarta. Data yang diambil merupakan postingan Twitter yang berasal dari hasil pencarian dengan kata kunci “Pilpres 2019”. Pengambilan data dengan cara *crawling* yang memanfaatkan fasilitas API Twitter. Metode yang digunakan adalah *extended (full text)* dengan *library tweepy* dan tanpa menggunakan *retweet*. Data yang terkumpul berjumlah 5055 yang merupakan gabungan data *tweet* dengan kata kunci “Pilpres 2019” dan yang berhubungan dengan hal tersebut.

2.2 Seleksi, pelabelan data dan *preprocessing*

Berdasarkan pendekatan eksperimental dalam penelitian ini, dalam proses seleksi diambil 3216 data latih dan data sejumlah 5055 sebagai data uji. Proses seleksi ini hanya dilakukan pada data yang akan digunakan dalam proses pembelajaran. Pemberian label dilakukan secara manual berdasarkan pada definisi dari pakar komunikasi. Setiap postingan *tweet* yang sudah diambil akan dikelompokkan ke dalam empat kelas yaitu positif (1), negatif (-1), netral (0), serta irrelevant (2). Pada tahap *preprocessing* dilakukan proses pembersihan terhadap data *tweet*. Langkah-langkah dalam pembersihan datanya meliputi : *cleansing*, *case folding*, *tokenizing*, *convert slangword*, *stopword removal* dan *stemming*. Semua proses ini dilakukan dengan tujuan agar kualitas klasifikasi yang dilakukan menjadi lebih baik [10].

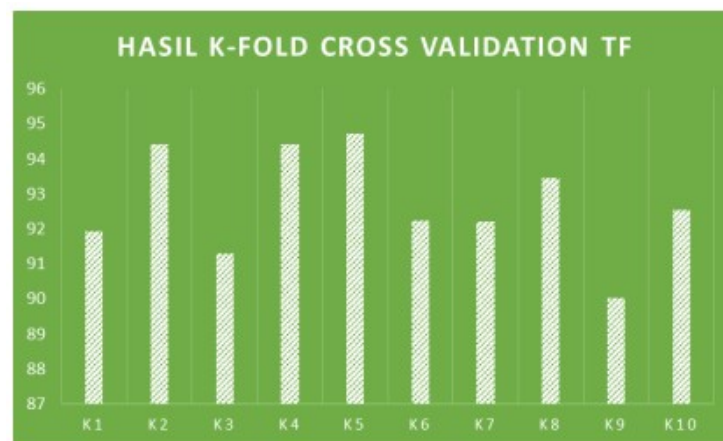
3 Hasil dan pembahasan

Data latih yang diambil sebanyak 3216 data yang telah melalui *preprocessing* dan diberi label secara manual akan dilakukan ekstraksi fitur. Ekstraksi fitur ini digunakan untuk mencari *Term Frequency* (TF) dan *Term Frequency-Inverse Document Frequency* (TF-IDF). Ekstraksi fitur ini digunakan dalam perhitungan probabilitas menggunakan NBC sebagai metode klasifikasi. Setelah melalui proses klasifikasi, maka diperlukan penilaian tingkat akurasi dari proses klasifikasi yang sudah dilakukan. Evaluasi yang dilakukan dengan menggunakan metode *K-Fold Cross validation* yang berguna untuk mengetahui rata-rata keberhasilan dari suatu sistem dengan cara melakukan perulangan dengan mengacak atribut masukan sehingga sistem tersebut teruji untuk beberapa atribut input yang acak [11].

Fold yang digunakan adalah 10-fold, yaitu membagi data latih menjadi sepuluh bagian. *K-fold cross validation* pada iterasi k-1 menjadikan 1-300 data pertama menjadi data uji, sedangkan data ke 301-3000 menjadi data training. Setelah semua iterasi dilakukan maka didapatkan rata-rata akurasi dari klasifikasi sentimen. Untuk penggunaan jumlah fold terbaik sebagai uji validitas, dianjurkan menggunakan 10-fold *cross validation* dalam model [12].

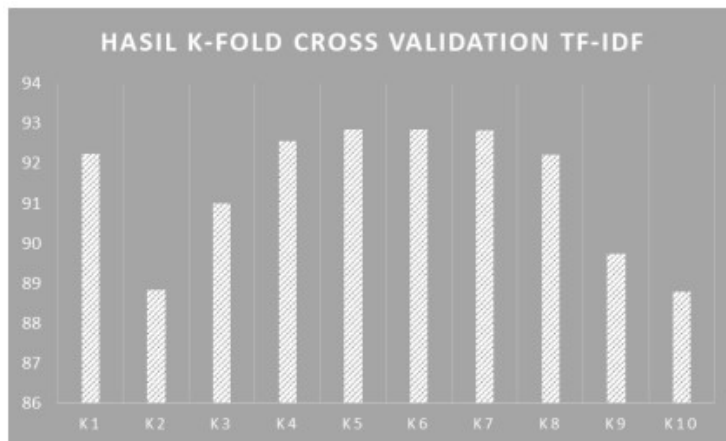
Rata-rata akurasi analisis sentimen yang dihasilkan dengan menggunakan metode NBC dan pembobotan TF pada data latih tweet ujaran kebencian sebesar 92,7%. Sedangkan rata-rata akurasi sentimen menggunakan metode NBC dan pembobotan TF-IDF pada data latih *tweet* ujaran kebencian sebesar 91,3%. Hal ini menunjukkan bahwa klasifikasi *multinomial* Naive Bayes menggunakan nilai TF dalam perhitungan nilai peluang frekuensi kata dari setiap kelas bekerja dengan baik, dimana nilai peluang frekuensi kata tersebut disimpan dalam tabel naive bayes yang menyimpan nilai peluang frekuensi kata dalam data latih dan term atau kata.

Data uji dilakukan pengecekan kata dalam daftar data latih, apabila kata dalam data uji tidak ada dalam daftar data latih maka akan diabaikan [13]. Hasil perhitungan peluang data uji terhadap data latih dilihat besarnya, semakin besar nilai peluangnya maka semakin tinggi tingkat peluang data uji masuk dalam kelas tersebut. Perbandingan tingkat akurasi ditunjukkan pada gambar 2 dan 3.



Gambar 2 Grafik akurasi K-fold cross validation TF

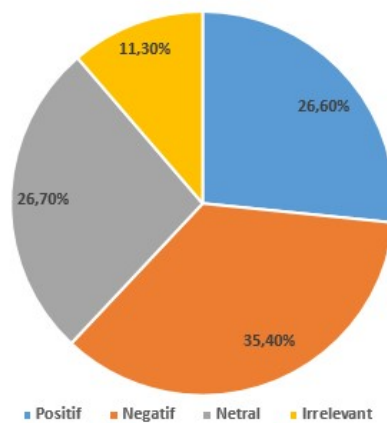
Proses berikutnya adalah melakukan klasifikasi pada data uji berdasarkan model klasifikasi dari data latih yang telah diolah sebelumnya dengan menggunakan NBC dan pembobotan



■ **Gambar 3** Grafik akurasi K-fold cross validation TF-IDF

TF-IDF. Data uji ini berjumlah 5055 data *tweet*. Setelah dilabeli secara otomatis kemudian didapatkan hasil klasifikasi sentimen irrelevant sebanyak 11,3% dengan jumlah data 573, sentimen negatif sebanyak 35,4% dengan jumlah data 1786, sentimen netral sebanyak 26,7% dengan jumlah data 1350 dan sentimen positif sebanyak 26,6% dengan jumlah data 1343 seperti yang ditunjukkan pada Gambar 4.

**Klasifikasi Analisis Sentimen : Ujaran Kebencian
#pilpres2019**

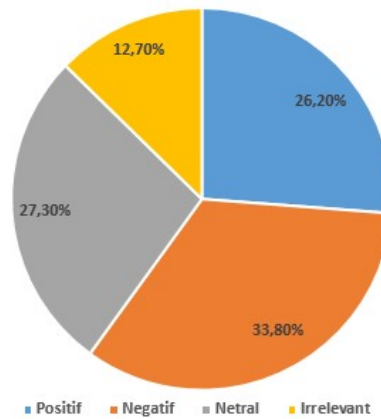


■ **Gambar 4** Ujaran kebencian #pilpres2019

Hasil klasifikasi ini kemudian dipilah-pilah berdasarkan sentimen analisis di beberapa kota besar di Indonesia. Kota Jakarta memperoleh hasil kecenderungan sentimen negatif yang paling dominan diantara yang lain yaitu sejumlah 33,8%, dan diikuti sentimen netral sejumlah 27,3%, positif sejumlah 26,2% dan irrelevant sejumlah 12,7%. Sehingga dapat disimpulkan bahwa sentimen ujaran kebencian di kota Jakarta yang paling besar adalah sentimen negatif yaitu dengan nilai presentase 33,8%. Gambar 5 memperlihatkan sebaran

ujaran kebencian di kota Jakarta.

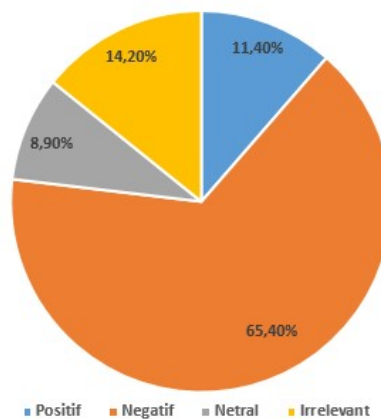
**Klasifikasi Analisis Sentimen : Ujaran Kebencian
#pilpres2019 di kota Jakarta**



Gambar 5 Ujaran Kebencian #pilpres2019 di kota Jakarta

Kota Bandung memperoleh hasil kecenderungan sentimen negatif yang paling dominan diantara yang lain yaitu sejumlah 65,4%, irrelevant sejumlah 14,2%. positif sejumlah 11,4% dan sentimen netral sejumlah 8,9 %. Sehingga dapat disimpulkan bahwa sentimen ujaran kebencian di kota Bandung yang paling besar adalah sentimen negatif yaitu dengan nilai presentase 65,4% seperti yang ditunjukkan dalam Gambar 6.

**Klasifikasi Analisis Sentimen : Ujaran Kebencian
#pilpres2019 di kota Bandung**

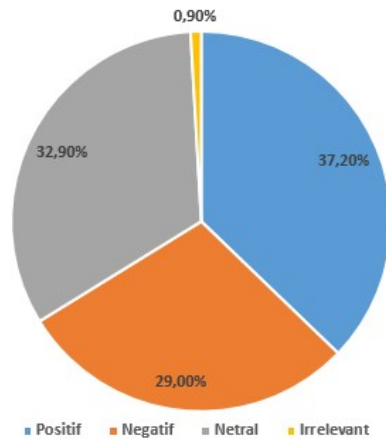


Gambar 6 Ujaran kebencian #pilpres2019 di kota Bandung

Kota Surabaya memperoleh hasil kecenderungan sentimen positif yang paling dominan yaitu 37,2 %, sentimen netral sejumlah 32,9%, negatif sejumlah 29,0% dan irrelevant sejumlah 0,9%. Sehingga dapat disimpulkan bahwa sentimen ujaran kebencian di kota Surabaya yang

paling besar adalah sentimen positif yaitu dengan nilai presentase 37,2 %. Gambar 7 menunjukkan sebaran sentimen di Kota Surabaya. Kota Yogyakarta, yang memperoleh hasil

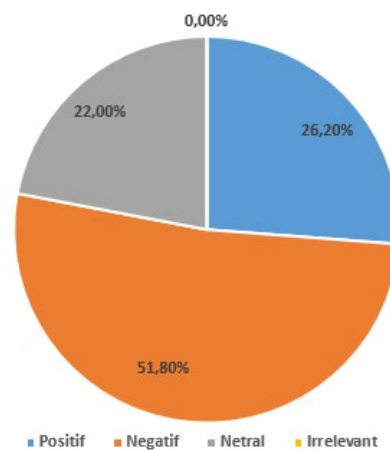
**Klasifikasi Analisis Sentimen : Ujaran Kebencian
#pilpres2019 di kota Surabaya**



Gambar 7 Ujaran kebencian #pilpres2019 di kota Surabaya

kecenderungan sentimen negatif paling dominan diantara yang lain yaitu sejumlah 51,8%, dan diikuti sentimen positif sejumlah 26,2%, netral sejumlah 22,0% dan irrelevant sejumlah 0%. Sehingga dapat disimpulkan bahwa sentimen ujaran kebencian di kota Yogyakarta yang paling besar adalah sentimen negatif yaitu dengan nilai presentase 51,8%, seperti yang diperlihatkan dalam Gambar 8

**Klasifikasi Analisis Sentimen : Ujaran Kebencian
#pilpres2019 di kota Yogyakarta**

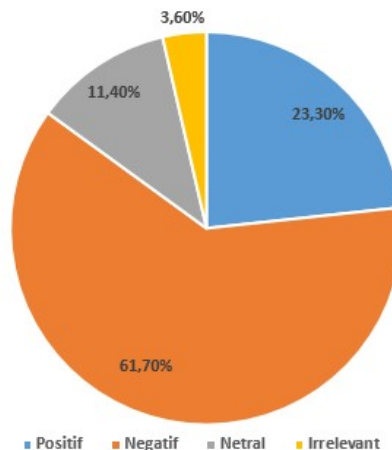


Gambar 8 Ujaran kebencian #pilpres2019 di kota Yogyakarta

Kota Semarang memperoleh hasil kecenderungan sentimen negatif yang paling dominan diantara yang lain yaitu sejumlah 61,7%, dan diikuti sentimen positif sejumlah 23,3%, netral

sejumlah 11,4% dan irrelevant sejumlah 3,6%. Sehingga dapat disimpulkan bahwa sentimen ujaran kebencian 63 pada pilpres 2019 di kota Semarang yang paling besar adalah sentimen negatif yaitu dengan nilai presentase 61,7% (Gambar 9).

**Klasifikasi Analisis Sentimen : Ujaran Kebencian
#pilpres2019 di kota Semarang**



Gambar 9 Ujaran kebencian #pilpres2019 di kota Semarang

4 Kesimpulan dan saran

Analisis sentimen dengan metode NBC dapat digunakan untuk mengklasifikasi ujaran kebencian dalam *tweet* dengan *hashtag* pilpres 2019. Hasil akurasi yang lebih baik didapatkan dengan menggunakan pembobotan TF yaitu 92,7% dibandingkan dengan pembobotan TF-IDF yaitu 91,3%. Akurasi sebesar 92,7% dan 91,3% ini merupakan rata-rata akurasi dari evaluasi model klasifikasi menggunakan k-fold cross validation pada 3216 data latih.

Hasil analisis sentimen ujaran kebencian pada data uji sejumlah 5055 data *tweet* dengan *hashtag* pilpres2019 di kota Jakarta, Bandung, Semarang, Surabaya, Yogyakarta dengan memanfaatkan model klasifikasi dari data latih menggunakan NBC dan Pembobotan TF-IDF didapatkan hasil klasifikasi sentimen irrelevant sebanyak 11,3% dengan 573 data, sentimen negatif sebanyak 35,4% dengan 1786 data, sentimen netral sebanyak 26,7% sebanyak 1350 data dan sentimen positif sebanyak 26,6% sebanyak 1343 data. Dengan kecenderungan pada sentimen negatif dengan nilai terbesar yaitu 35,4% di lima kota tersebut 17 3. Sedangkan hasil sentimen ujaran kebencian yang paling besar pada masing-masing kota yaitu: Jakarta dengan sentimen negatif sebesar 33,8%, Bandung dengan sentimen negatif sebesar 65,4%, Surabaya dengan sentimen positif sebesar 37,2 %,Yogyakarta dengan sentimen negatif sejumlah 51,8%. dan Semarang dengan sentimen negatif 61,7%.

Tidak bisa dipungkiri bahwa pembentukan opini di Twitter banyak dilakukan oleh tim sukses masing-masing kandidat dengan berbagai tujuan. Sehingga perlu dilakukan pengamatan lebih dalam apakah postingan-postingan tersebut dilakukan oleh akun-akun tertentu untuk tujuan kampanye hitam menjatuhkan masing-masing kandidat. Hal ini dapat dilakukan agar penarikan kesimpulan lebih bisa menunjukkan apa yang sedang terjadi.

Pustaka

- 1 A. Agarwal, B. Xie, I. Vovsha, O. Rambow, and R. J. Passonneau, "Sentiment analysis of twitter data," in *Proceedings of the Workshop on Language in Social Media (LSM 2011)*, 2011, pp. 30–38.
- 2 D. Farber, "Twitter hits 400 million tweets per day, mostly mobile." [Online]. Available: <https://www.cnet.com/news/twitter-hits-400-million-tweets-per-day-mostly-mobile/>
- 3 A. Pak and P. Paroubek, "Twitter as a corpus for sentiment analysis and opinion mining." in *LREc*, vol. 10, no. 2010, 2010, pp. 1320–1326.
- 4 Semiocast, "Geolocation analysis of twitter accounts and tweets by semiocast." [Online]. Available: https://semiocast.com/en/publications/2012_07_30_Twitter_reaches_half_a_billion_accounts_140m_in_the_US
- 5 F. Handayani and F. S. Pribadi, "Implementasi algoritma naive bayes classifier dalam pengklasifikasian teks otomatis pengaduan dan pelaporan masyarakat melalui layanan call center 110," *Jurnal Teknik Elektro*, vol. 7, no. 1, pp. 19–24, 2015.
- 6 S. Natalius, "Metoda naïve bayes classifier dan penggunaannya pada klasifikasi dokumen," *Makalah II2092 Probabilitas dan Statistik – Sem. I Tahun 2010/2011*, vol. 2010, 2010.
- 7 S. L. B. Ginting and R. P. Trinanda, "Teknik data mining menggunakan metode bayes classifier untuk optimalisasi pencarian pada aplikasi perpustakaan (studi kasus: Perpustakaan universitas pasundan–bandung)," *Jurnal Teknologi dan Informasi*, vol. 3, no. 2, pp. 37–50, 2013.
- 8 M. S. Mustafa, M. R. Ramadhan, and A. P. Thenata, "Implementasi data mining untuk evaluasi kinerja akademik mahasiswa menggunakan algoritma naive bayes classifier," *Creative Information Technology Journal*, vol. 4, no. 2, pp. 151–162, 2018.
- 9 A. F. Hidayatullah, A. S. Azhari *et al.*, "Analisis sentimen dan klasifikasi kategori terhadap tokoh publik pada twitter," in *Seminar Nasional Informatika 2014*. " Veteran" University of National Development Yogyakarta.
- 10 A. F. Hidayatullah, "Pengaruh stopword terhadap performa klasifikasi tweet berbahasa indonesia," *JISKA (Jurnal Informatika Sunan Kalijaga)*, vol. 1, no. 1, 2016.
- 11 P. Pitria, "Analisis sentimen pengguna twitter pada akun resmi samsung indonesia dengan menggunakan naïve bayes," Ph.D. dissertation, Universitas Komputer Indonesia, 2014.
- 12 R. Kohavi *et al.*, "A study of cross-validation and bootstrap for accuracy estimation and model selection," in *Ijcai*, vol. 14, no. 2. Montreal, Canada, 1995, pp. 1137–1145.
- 13 D. Jufarsky, J. H. Martin, D. Jufarsky, and J. H. Martin, "Speech and language processing: an introduction to natural language processing, computational linguistics, and speech recognition," 2000.