

TESIS

ANALISIS SENTIMEN *TWEET* PUBLIK MENGGUNAKAN METODE *HETEROGENEOUS* DAN *HOMOGENEOUS* *ENSEMBLE LEARNING* TERHADAP PEMBANGUNAN IBU KOTA NEGARA INDONESIA

*Sentiment Analysis of Public Tweets Using Heterogeneous and
Homogeneous Ensemble Learning Methods on the Development of
Indonesia's National Capital City*

**MAWARDI KUDIN
D082211002**



**PROGRAM STUDI MAGISTER TEKNIK INFORMATIKA
DEPARTEMEN TEKNIK INFORMATIKA
FAKULTAS TEKNIK
UNIVERSITAS HASANUDDIN
GOWA
2024**



PENGAJUAN TESIS

ANALISIS SENTIMEN TWEET PUBLIK MENGGUNAKAN METODE HETEROGENEOUS DAN HOMOGENEOUS ENSEMBLE LEARNING TERHADAP PEMBANGUNAN IBU KOTA NEGARA INDONESIA

Tesis

Sebagai Salah Satu Syarat untuk Mencapai Gelar Magister
Program Studi Teknik Informatika

Disusun dan diajukan oleh

**MAWARDI KUDIN
D082211002**

Kepada

**FAKULTAS TEKNIK
UNIVERSITAS HASANUDDIN
GOWA
2024**



TESIS

ANALISIS SENTIMEN TWEET PUBLIK MENGGUNAKAN METODE HETEROGENEOUS DAN HOMOGENEOUS ENSEMBLE LEARNING TERHADAP PEMBANGUNAN IBU KOTA NEGARA INDONESIA

MAWARDI KUDIN
D082211002

Telah dipertahankan di hadapan Panitia Ujian Tesis yang dibentuk dalam rangka penyelesaian studi pada
Program Magister Teknik Informatika Fakultas Teknik
Universitas Hasanuddin
Pada tanggal 08 Maret 2024
dan dinyatakan telah memenuhi syarat kelulusan

Menyetujui,

Pembimbing Utama



Dr. Amil Ahmad Ilham, S.T., M.IT.
NIP. 19731010 199802 1 001

Pembimbing Pendamping



Dr. Eng. Ady Wahyudi Paundu, S.T., M.T.
NIP. 19750313 200912 1 003

Dekan Fakultas Teknik
Universitas Hasanuddin



nad Isran Ramli, M.T. IPM., ASEAN.Eng.
9730926 200012 1 002

Ketua Program Studi
S2 Teknik Informatika



Dr. Ir. Zahir Zainuddin, M.Sc.
NIP. 19640427 198910 1 002



P1

PERNYATAAN KEASLIAN TESIS DAN PELIMPAHAN HAK CIPTA

Yang bertanda tangan di bawah ini

Nama : Mawardi Kudin
Nomor mahasiswa : D082211001
Program studi : S2 Teknik Informatika

Dengan ini menyatakan bahwa, tesis berjudul “Analisis Sentimen Tweet Publik Menggunakan Metode Heterogeneous dan Homogeneous Ensemble Learning Terhadap Pembangunan Ibu Kota Negara Indonesia” adalah benar karya saya dengan arahan dari komisi pembimbing Dr. Amil Ahmad Ilham, S.T., M.IT. dan Dr. Eng. Ady Wahyudi Paundu, S.T., M.T. Karya ilmiah ini belum diajukan dan tidak sedang diajukan dalam bentuk apa pun kepada perguruan tinggi mana pun. Sumber informasi yang berasal atau dikutip dari karya yang diterbitkan maupun tidak diterbitkan dari penulis lain telah disebutkan dalam teks dan dicantumkan dalam Daftar Pustaka tesis ini. Sebagian dari isi tesis ini telah dipublikasikan di Prosiding 2023 IEEE International Conference on Communication, Networks and Satellite (COMNETSAT), halaman 164-169, dan DOI: 10.1109/COMNETSAT59769.2023.10420690 sebagai artikel dengan judul “Performance Enhancement of Individual Learning Methods for Sentiment Analysis Using Ensemble Learning and Soft Voting Techniques”.

Dengan ini saya melimpahkan hak cipta dari karya tulis saya berupa tesis ini kepada Universitas Hasanuddin.

Gowa, 04 April 2024

Yang menyatakan


Mawardi Kudin



KATA PENGANTAR

Puji Syukur penulis panjatkan kehadirat Allah SWT, berkat rahmat, petunjuk, bimbingan dan karunia-Nya sehingga Karya Ilmiah berupa Tesis dengan judul “**Analisis Sentimen Tweet Publik Menggunakan Metode *Heterogeneous* dan *Homogeneous Ensemble Learning* Terhadap Pembangunan Ibu Kota Negara Indonesia**” ini dapat selesai. Tesis ini diajukan untuk memenuhi salah satu syarat dalam mendapatkan gelar Magister Komputer (M. Kom.) pada Program Studi Magister Teknik Informatika, Departemen Teknik Informatika, Fakultas Teknik, Universitas Hasanuddin.

Penulis menyadari bahwa selama penulisan dan penyusunan tesis ini banyak pihak yang telah membantu dan memberikan dukungannya baik secara material maupun moril kepada penulis. Demikian pula segala bantuan yang penulis peroleh selama di bangku perkuliahan sehingga dapat membantu penulis dalam menyusun dan menyelesaikan tesis ini.

Oleh karena itu Penulis menyampaikan ucapan terima kasih yang tulus kepada keluarga penulis, terkhusus kepada kedua orang tua, ayahandaku Muh. Kudin dan ibundaku Hasnah serta kedua kakakku yang tercinta Muliani Kudin dan Muis Kudin yang menjadi alasan penulis untuk tetap bertahan hingga saat ini. Tidak lupa pula penulis ucapkan terima kasih yang setulusnya dan penghormatan kepada bapak Dr. Amil Ahmad Ilham, S.T., M.IT. selaku Pembimbing I dan bapak Dr. Eng. Ady Wahyudi Paundu, S.T., M.T. selaku Pembimbing II, yang telah meluangkan waktu, tenaga, dan pikiran serta memberikan arahan-arahan yang sangat berharga bagi penulis, sehingga penulis dapat menyelesaikan penelitian ini.

Selanjutnya, penghargaan dan rasa terima kasih yang setinggi-tingginya juga penulis haturkan kepada:

1. Bapak Dr. Ir. Zahir Zainuddin, M.Sc. selaku Ketua Prodi Magister Teknik Informatika yang selalu memberikan arahan kepada Penulis agar cepat menyelesaikan segala rangkaian proses akademik di Departemen Teknik Informatika.

2. Prof. Dr. Ir. Indrabayu, ST., MT., M.Bus.Sys., IPM, ASEAN. Eng.

3. Penguji I, Bapak Dr. Ir. Zahir Zainuddin, M.Sc. selaku Penguji II, dan

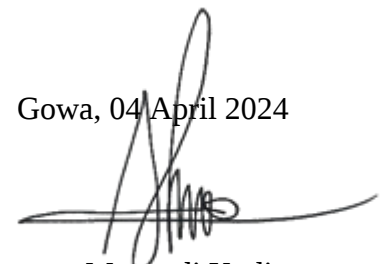


Bapak Prof. Dr. Ir. H. Ansar Suyuti, MT., SH., IPU., ASEAN.Eng. selaku Penguji III yang senantiasa meluangkan waktunya untuk memberikan kritik serta saran dalam pembuatan tesis ini.

3. Bapak dan Ibu Dosen di Departemen Teknik Informatika yang telah memberikan ilmunya selama proses perkuliahan.
4. Seluruh staff Program Studi Magister Teknik Informatika yang telah memberikan bantuan administrasi selama penulis berkuliah di Fakultas Teknik, Universitas Hasanuddin.
5. Teman-teman mahasiswa angkatan 4 serta seluruh pengurus dan anggota Himpunan Mahasiswa Magister Teknik Informatika yang senantiasa memberikan dukungan dan semangat kepada Penulis dalam menyelesaikan penelitian.
6. Serta seluruh pihak yang tidak bisa disebutkan satu persatu, terima kasih atas semua dukungan yang telah diberikan dalam penyelesaian tesis ini.

Penulis menyadari masih banyak kekurangan yang terdapat pada tesis ini. Saran dan kritik penulis harapkan dari pembaca untuk perbaikan-perbaikan dimasa yang akan datang.

Gowa, 04 April 2024



Mawardi Kudin



ABSTRAK

MAWARDI KUDIN. Analisis Sentimen *Tweet* Publik Menggunakan Metode *Heterogeneous* dan *Homogeneous Ensemble Learning* Terhadap Pembangunan Ibu Kota Negara Indonesia (dibimbing oleh **Amil Ahmad Ilham** dan **Ady Wahyudi Paundu**).

Pemerintah memulai program pemindahan Ibu Kota Negara (IKN) ke Kalimantan yang menuai pro dan kontra. Isu ini mendapatkan perbincangan hangat di masyarakat dan media sosial. Pembangunan IKN menghadapi berbagai tanggapan atau pendapat dari masyarakat Indonesia, mulai dari positif, netral, hingga negatif. Namun, sulit untuk mengetahui pendapat mayoritas masyarakat secara langsung karena akan memakan banyak tenaga dan waktu yang sangat lama. Pesatnya perkembangan teknologi saat ini meningkatkan pemanfaatan media sosial, khususnya Twitter, sebagai tempat untuk mengekspresikan pendapat. Kita dapat mengetahui opini seseorang berdasarkan setiap *tweet* yang ditulisnya menggunakan teknik Analisis Sentimen, yaitu sebuah metode untuk menentukan apakah sentimen dalam suatu teks itu positif, netral, atau negatif. Dalam melakukan analisis sentimen, telah banyak diterapkan metode *individual learning*, yaitu hanya menggunakan sebuah algoritma *machine learning* saja dalam melakukan klasifikasi sentimen. Namun, metode ini akan memiliki kekurangan tergantung setiap algoritma yang digunakan. Untuk menyelesaikan masalah tersebut, telah dikembangkan metode *ensemble learning* yang menggabungkan beberapa algoritma *machine learning* untuk mendapatkan performa yang lebih baik. Dalam penelitian ini diusulkan penerapan metode *heterogeneous* dan *homogeneous ensemble learning* dalam melakukan analisis sentimen, dengan menggabungkan beberapa algoritma *machine learning* yang sejenis maupun yang berbeda, dan metode *weighted soft voting* diterapkan untuk menentukan hasil akhir klasifikasi sentimen pada setiap metode *ensemble learning* yang diusulkan. Selain itu, metode TF-IDF juga digunakan untuk vektorisasi data yang berbentuk teks. Skor akurasi tertinggi dalam melakukan analisis sentimen diperoleh oleh metode *heterogeneous ensemble learning* yang diusulkan, yaitu sebesar 76%. Skor ini berhasil mengungguli akurasi dari metode *individual learning*.

Kata Kunci: Analisis Sentimen, *Ensemble Learning*, *Heterogeneous*, *Homogeneous*, *Weighted Soft Voting*, TF-IDF, Twitter.



ABSTRACT

MAWARDI KUDIN. Sentiment Analysis of Public Tweets Using Heterogeneous and Homogeneous Ensemble Learning Methods on the Development of Indonesia's National Capital City (supervised by **Amil Ahmad Ilham** and **Ady Wahyudi Paundu**).

The government has initiated a program to relocate the National Capital City (NCC) to Kalimantan, which has sparked both support and opposition. This issue has been widely discussed in society and on social media. The development of the NCC has generated various responses or opinions from the Indonesian public, ranging from positive to neutral and negative. However, directly determining the majority opinion of the public would require a significant amount of time and effort. The rapid advancement of technology has led to increased use of social media, particularly Twitter, as a platform for expressing opinions. By employing Sentiment Analysis, a technique used to determine whether the sentiment in a text is positive, neutral, or negative, we can gain insights into individuals' opinions based on their tweets. In conducting sentiment analysis, the widely used approach is individual learning, where only a single machine learning algorithm is used for sentiment classification. However, this method has its limitations depending on the specific algorithm employed. An ensemble learning method has been developed to address this issue, combining multiple machine learning algorithms to improve performance. This study proposes implementing heterogeneous and homogeneous ensemble learning methods for sentiment analysis. These methods involve combining multiple machine learning algorithms of the same or different types. The weighted soft voting method determines the final sentiment classification result for each ensemble learning method proposed. Additionally, the TF-IDF method is used for text data vectorization. The proposed heterogeneous ensemble learning method achieves the highest accuracy score in sentiment analysis, reaching 76%. This score outperforms the accuracy obtained by the individual learning method.

Keywords: Sentiment Analysis, Ensemble Learning, Heterogeneous, Homogeneous, Weighted Soft Voting, TF-IDF, Twitter.



DAFTAR ISI

HALAMAN JUDUL.....	i
PENGAJUAN TESIS	ii
PERSETUJUAN TESIS	iii
PERNYATAAN KEASLIAN TESIS	iv
KATA PENGANTAR	v
ABSTRAK	vii
ABSTRACT	viii
DAFTAR ISI.....	ix
DAFTAR TABEL	xii
DAFTAR GAMBAR	xiii
DAFTAR LAMPIRAN	xiv
BAB I PENDAHULUAN	1
1.1 Latar Belakang	1
1.2 Rumusan Masalah	6
1.3 Tujuan Penelitian.....	7
1.4 Manfaat Penelitian.....	7
1.5 Batasan Masalah.....	7
BAB II TINJAUAN PUSTAKA.....	8
2.1 Kajian Pustaka	8
2.1.1 Analisis sentimen.....	8
2.1.2 <i>Ensemble learning</i>	12
2.1.3 <i>Naïve bayes</i>	15
2.1.4 SVM (<i>Super Vector Machine</i>).....	15
2.1.5 <i>Logistic regression</i>	16
2.1.6 <i>Heterogeneous dan homogeneous ensemble learning</i>	16
2.1.7 <i>Weighted soft voting</i>	19
2.1.8 <i>Bagging (Bootstrap Aggregating)</i>	23
2.2 <i>TF-IDF (Term Frequency - Inverse Document Frequency)</i>	24
2.3 <i>0 Metaheuristic dan PSO (Particle Swarm Optimization)</i>	26
2.4 Metode Penyelesaian Masalah	27



2.2.1 <i>State of the art</i> penelitian	27
2.2.2 Metode yang diusulkan	34
2.3 Target Hipotesis Hasil Penelitian	36
2.4 Kerangka Pikir Penelitian	37
BAB III METODOLOGI PENELITIAN	38
3.1 Jenis Penelitian	38
3.2 Tahapan Penelitian	38
3.3 Sumber Data	39
3.4 Rancangan Sistem	40
3.4.1 Tahapan pra-pemrosesan	40
3.4.2 <i>Ensemble learning classification</i>	41
3.5 Proses Kerja Vektorisasi Data Text	47
3.6 Proses Kerja Setiap Algoritma yang Digunakan	49
3.7 Pengujian Sistem	58
3.8 Waktu dan Lokasi Penelitian	59
BAB IV HASIL DAN PEMBAHASAN	60
4.1 Hasil Pengumpulan dan Pelabelan <i>Dataset</i>	60
4.2 <i>Preprocessing</i>	61
4.3 Pembagian <i>Dataset</i> Menjadi Data Latih dan Data Uji	64
4.4 Vektorisasi Data	65
4.5 Evaluasi Kinerja Tiga Metode <i>Individual Learning</i>	67
4.6 Evaluasi Hasil <i>Weighted Soft Voting</i>	68
4.6.1 Probabilitas untuk setiap kelas	68
4.6.2 Bobot untuk setiap model	69
4.7 Evaluasi Kinerja Metode <i>Heterogeneous</i> dan <i>Homogeneous Ensemble Learning</i> yang Diusulkan	72
4.7.1 Evaluasi kinerja metode <i>heterogeneous ensemble learning</i>	72
4.7.2 Evaluasi kinerja metode <i>homogeneous ensemble learning</i>	74
4.7.3 Perbandingan kinerja metode <i>heterogeneous</i> dan <i>homogeneous</i>	76
4.8 Evaluasi Perbandingan Kinerja dari Metode <i>Individual Learning</i> dan Metode <i>Ensemble Learning</i> yang Diusulkan	76
KESIMPULAN DAN SARAN	81



5.1 Kesimpulan.....	81
5.2 Saran.....	82
DAFTAR PUSTAKA	83
LAMPIRAN.....	88



DAFTAR TABEL

Tabel 1 <i>State of the art</i> penelitian	27
Tabel 2 <i>Confusion Matrix</i>	58
Tabel 3 Sampel <i>dataset</i> yang telah dikumpulkan	60
Tabel 4 Kinerja algoritma <i>Naïve Bayes</i> , <i>SVM</i> , dan <i>Logistic Regression</i>	68
Tabel 5 Hasil pengujian metode <i>Heterogeneous</i> dengan <i>weighted soft voting</i>	73
Tabel 6 Hasil pengujian metode <i>Heterogeneous</i> dengan <i>simple soft voting</i>	73
Tabel 7 Hasil pengujian metode <i>Homogeneous</i> dengan <i>weighted soft voting</i>	75
Tabel 8 Hasil pengujian metode <i>Homogeneous</i> dengan <i>simple soft voting</i>	75
Tabel 9 Perbandingan hasil pengujian metode <i>individual learning</i> dan <i>ensemble learning</i>	77



DAFTAR GAMBAR

Gambar 1 Tahapan analisis sentimen.....	9
Gambar 2 Tahapan <i>heterogeneous ensemble learning</i>	16
Gambar 3 Tahapan <i>homogeneous ensemble learning</i>	18
Gambar 4 Kerangka pikir penelitian	37
Gambar 5 Tahapan penelitian	38
Gambar 6 <i>Flowchart</i> proses analisis sentimen.....	40
Gambar 7 Tahapan pra-pemrosesan.....	41
Gambar 8 Tahapan klasifikasi <i>heterogeneous ensemble learning</i>	42
Gambar 9 Tahapan klasifikasi <i>homogeneous ensemble learning</i>	44
Gambar 10 Gambaran sederhana garis <i>hyperlane</i> yang memisahkan 2 kelas	53
Gambar 11 Gambaran sederhana penentuan <i>hyperplane</i> berdasarkan <i>margin</i>	54
Gambar 12 <i>Dataset</i> hasil proses <i>case folding</i>	61
Gambar 13 <i>Dataset</i> hasil proses <i>tokenization</i>	62
Gambar 14 <i>Dataset</i> hasil proses <i>filtering</i>	63
Gambar 15 <i>Dataset</i> hasil proses <i>stemming</i>	63
Gambar 16 <i>Dataset</i> setelah dibagi menjadi data latih dan data uji	64
Gambar 17 Jumlah setiap kelas pada data latih	65
Gambar 18 Jumlah setiap kelas pada data uji	65
Gambar 19 Hasil proses TF	66
Gambar 20 Hasil proses IDF.....	66
Gambar 21 Hasil proses TF-IDF.....	67
Gambar 22 Data setelah diubah kedalam bentuk <i>array</i>	67
Gambar 23 Nilai probabilitas setiap kelas dari hasil pengujian model NB	69
Gambar 24 Nilai probabilitas setiap kelas dari hasil pengujian model SVM	69
Gambar 25 Nilai probabilitas setiap kelas dari hasil pengujian model LR.....	69
Gambar 26 Kombinasi nilai bobot setiap model dari metode <i>heterogeneous</i>	70
Gambar 27 Kombinasi nilai bobot setiap model dari metode <i>homogeneous</i>	71
28 Nilai <i>Confusion Matrix</i> dari model <i>heterogeneous</i>	72
29 Nilai <i>Confusion Matrix</i> dari 3 model <i>homogeneous</i>	74
30 Grafik perbandingan metode <i>individual</i> dan <i>ensemble learning</i>	77



DAFTAR LAMPIRAN

Lampiran 1. <i>Script</i> Program untuk Prapemrosesan Data	89
Lampiran 2. <i>Script</i> Program untuk Vektorisasi Data menggunakan TF-IDF	91
Lampiran 3. <i>Script</i> Program Algoritma <i>Naïve Bayes</i>	92
Lampiran 4. <i>Script</i> Program Algoritma SVM.....	93
Lampiran 5. <i>Script</i> Program Algoritma <i>Logistic Regression</i>	94
Lampiran 6. <i>Script</i> Program <i>Confusion Matrix</i> untuk Menghitung Performa	95
Lampiran 7. <i>Script</i> Program <i>Heterogeneous Ensemble Learning</i>	96
Lampiran 8. <i>Script</i> Program <i>Homogeneous Ensemble Learning Naïve Bayes</i>	97
Lampiran 9. <i>Script</i> Program <i>Homogeneous Ensemble Learning SVM</i>	98
Lampiran 10. <i>Script</i> Program <i>Homogeneous Ensemble Learning LR</i>	99
Lampiran 11. Potongan <i>Script</i> Algoritma PSO untuk Model <i>Heterogeneous</i>	100
Lampiran 12. Potongan <i>Script</i> Algoritma PSO Model <i>Homogeneous NB</i>	101
Lampiran 13. Potongan <i>Script</i> Algoritma PSO Model <i>Homogeneous SVM</i>	102
Lampiran 14. Potongan <i>Script</i> Algoritma PSO Model <i>Homogeneous LR</i>	103



BAB I PENDAHULUAN

1.1 Latar Belakang

Isu pemindahan Ibu Kota Indonesia perlahan memanasi di media sosial setelah Presiden Joko Widodo menyetujui desain final bangunan Istana Negara tersebut. Nyoman Nuarta pun diundang ke Istana Merdeka untuk mempresentasikannya pada tanggal 3 Januari 2022 kemarin. Ibu Kota Negara ini dibangun untuk mencapai target Indonesia sebagai negara maju, sesuai Visi Indonesia 2045. Dibangun dengan identitas nasional, IKN akan mengubah orientasi pembangunan menjadi Indonesia-sentris, serta mempercepat Transformasi Ekonomi Indonesia.

Rapat Paripurna DPR yang dipimpin oleh Ketua DPR Puan Maharani telah menyetujui pengesahan Undang-Undang tentang Ibu Kota Negara pada tanggal 18 Januari 2022 lalu. Sebelumnya, UU IKN secara resmi mulai dibahas oleh Pansus IKN pada masa persidangan II tahun sidang 2021-2022 pada 7 Desember 2021. Pansus IKN melaksanakan rapat kerja dengan Menteri Perencanaan Pembangunan Nasional (PPN/Bappenas), Menteri Keuangan, Menteri ATR/BPN, Menteri Dalam Negeri, serta Menteri Hukum dan HAM. Bersamaan dengan rapat 18 Januari 2022 itulah turut disepakati juga pemberian nama “Nusantara” sebagai Ibu Kota Negara. Ketika melihat DPR telah mengesahkan RUU IKN dan memberikan nama “Nusantara” untuk Ibu Kota Negara Baru, tampaknya publik makin tersadar bahwa rencana pemindahan Ibu Kota bukan hanya sekadar wacana lagi.

Langkah pemerintah untuk memulai program pemindahan Ibu Kota Negara (IKN) ke Kalimantan ini ternyata masih menuai pro dan kontra. Isu ini terus berjalan dengan ragam topik yang mewarnainya, mulai dari pengesahan Undang-Undang Ibu Kota Negara (UU IKN), naskah akademik, hingga kondisi ekonomi. Kesadaran kolektif dari berbagai kalangan di masyarakat juga telah memunculkan perbincangan hangat tentang pemindahan Ibu Kota Negara, dan kini topik ini kian melebar di media sosial. Melihat berbagai isu dan permasalahan-permasalahan

terkait pembangunan Ibu Kota Negara Baru, tentunya akan memunculkan tanggapan dan kritikan dari masyarakat di Indonesia mengenai unan tersebut, mulai dari tanggapan positif, netral, hingga negatif.



Menganalisis dan mengetahui pendapat atau opini masyarakat mengenai pemindahan dan pembangunan Ibu Kota Negara ini akan dapat membantu melihat apakah upaya pemindahan Ibu Kota yang dilakukan oleh pemerintah ini telah mendapatkan dukungan dari mayoritas masyarakat Indonesia atau tidak. Akan tetapi, permasalahannya kita tidak mungkin menanyakan secara langsung ke seluruh masyarakat Indonesia tentang opini mereka mengenai pemindahan dan pembangunan Ibu Kota Negara, karena akan memakan banyak tenaga dan waktu yang sangat lama.

Perkembangan teknologi yang terus meningkat saat ini membuat penggunaan media sosial juga terus meningkat. Banyaknya pengguna media sosial saat ini membuat masyarakat biasanya menuangkan dan mengungkapkan berbagai pendapat atau sentimen mereka terhadap suatu hal melalui media sosial ini agar bisa dilihat dan dibagikan oleh masyarakat luas. Salah satu media sosial yang paling banyak digunakan untuk menuangkan pendapat adalah Twitter. Twitter adalah tempat di mana orang sering memilih untuk mengekspresikan emosi dan sentimen mereka tentang suatu merek, produk, atau layanan (V. Prakruthi et al. 2018) [1]. Apalagi jumlah pengguna Twitter telah mencapai 330 juta orang di seluruh dunia dan setiap detiknya menghasilkan 8000 data. *Tweet* yang disampaikan juga dapat berupa berita, opini, argumentasi, dan beberapa jenis kalimat lainnya (Kusrini and M. Mashuri 2019) [2]. Hal ini menyebabkan Twitter kaya akan teks yang memiliki data tertentu dan sangat cocok sebagai tempat untuk melihat bagaimana pandangan atau pendapat masyarakat terhadap suatu hal.

Memanfaatkan media sosial Twitter untuk mengambil dan menganalisis setiap pendapat atau sentimen pada setiap *tweet*, dapat membantu dalam melihat dan menentukan sentimen setiap orang apakah itu positif, netral, atau negatif terhadap suatu hal, terkhususnya sentimen masyarakat Indonesia terhadap pembangunan Ibu Kota Negara Baru “Nusantara”. Untuk mendapatkan jenis sentimen dari setiap *tweet* yang berisikan pendapat masyarakat, akan digunakan teknik Analisis Sentimen. Analisis Sentimen adalah metode untuk menentukan apakah sentimen

itu positif, netral, atau negatif. Ini juga dikenal sebagai polaritas material pembangunan opini (L. Mandloi and R. Patel 2020) [3]. Analisis sentimen dilakukan pengguna untuk mengetahui keadaan ulasan yang dibuat mengenai



produk tertentu atau hal-hal yang sedang berlangsung (V. Prakruthi et al. 2018) [1]. Tujuan dari analisis sentimen ini adalah untuk mengekstrak informasi dari komentar di media sosial dalam bentuk data teks, data ini kemudian dapat diolah menjadi informasi dan pengetahuan yang berguna tentang masalah yang berkembang (F. A. Wenando et al. 2020) [4], sehingga data dari komentar tersebut dapat diklasifikasikan ke dalam sentimen positif, netral, atau bahkan negatif.

Dalam melakukan analisis sentimen, sebelum *tweet* diklasifikasikan jenis sentimennya, ada beberapa tahapan yang harus dilakukan terlebih dahulu, mulai dari pengumpulan data *tweet* menggunakan Twitter API, hingga melakukan pra-pemrosesan terhadap data *tweet* yang masih berantakan. Menurut penelitian yang dilakukan D. A. Damaratih (2021) [5], ada beberapa teknik yang dapat digunakan dalam tahapan pra-pemrosesan diantaranya adalah *case folding*, *tokenization*, *stopword removal/filtering*, dan *stemming*. Selain itu, karena model *machine learning* tidak dapat memahami data tekstual, jadi kita harus memasukkan nilai numerik ke model machine learning. Untuk melakukannya, kita harus mengubah data tekstual ke dalam bentuk vektor menggunakan teknik *Natural Language Processing* untuk vektorisasi agar dapat diproses selanjutnya. Ada dua teknik *Natural Language Processing* yang populer untuk vektorisasi, yaitu *Bag-of-Words* (BoW) dan *Term Frequency-Inverse Document Frequency* (TF-IDF) (M. T. H. K. Tusar and M. T. Islam 2021) [6]. Setelah melalui beberapa tahapan tersebut, barulah data *tweet* yang bersih akan diklasifikasikan ke dalam sentimen positif, netral, atau negatif menggunakan beberapa algoritma klasifikasi *machine learning*.

Dalam melakukan proses klasifikasi data, banyak algoritma atau metode yang dapat digunakan untuk membuat sebuah model klasifikasi, seperti *naïve bayes*, *Super Vector Machine* (SVM), *decision tree*, KNN, *logistic regression*, *neural network*, dan lain sebagainya. Penerapan setiap algoritma tersebut dalam proses pengolahan data seperti klasifikasi biasa juga disebut dengan *individual learning*, yaitu penggunaan sebuah algoritma (*individual algorithm*) untuk membuat sebuah model *machine learning* yang selanjutnya model tersebut dapat langsung digunakan



lakukan pengolahan data.

gunaan metode *individual learning* ini memang cukup efektif dalam gun model klasifikasi, tapi metode ini juga memiliki beberapa kekurangan,

seperti pada algoritma SVM, *decision tree* dan *logistic regression* akan sulit digunakan bila dipakai pada data dengan skala yang besar dan memiliki banyak fitur. Kemudian untuk *naïve bayes* memiliki kekurangan yaitu *independence* antar atribut membuat akurasi menjadi berkurang dan apabila probabilitas kondisional bernilai nol, maka probabilitas prediksi juga akan bernilai nol. Selanjutnya, untuk KNN memiliki waktu pemrosesan yang lama jika *dataset*-nya sangat besar dan sangat sensitif pada data yang memiliki banyak *noise*. Begitu pun juga untuk *neural network*, memiliki kekurangan dengan sifatnya yang seperti *black box*, dimana kita sulit melihat apa yang sebenarnya terjadi di dalam proses pembuatan model dan tidak mudah menemukan fitur apa yang penting bagi akurasi model. Setiap algoritma pasti memiliki kekurangannya sendiri, sehingga perlunya pengembangan pada metode *individual learning* ini agar bisa menutupi kekurangan setiap algoritma yang ada.

Menjawab dari pernyataan di atas, telah dikembangkan sebuah metode yang bernama *ensemble learning*. Metode *ensemble learning* ini merupakan metode yang menggabungkan beberapa algoritma *machine learning* untuk mendapatkan performa yang lebih baik, yang dimana ide dasar dari *ensemble learning* ini adalah menerapkan model pembelajaran yang berbeda-beda untuk mendapatkan klasifikasi atau hasil yang lebih baik dari *individual learning*. Hasil dari model klasifikasi yang berbeda dapat digabungkan dengan dua cara yang berbeda pula, yaitu *voting* dan *averaging* (P. S. Mung and S. Phyu 2020) [7]. Metode *ensemble learning* ini lebih efektif dibandingkan dengan metode *individual learning*, hasil ini telah dibuktikan oleh beberapa penelitian, seperti pada penelitian yang dilakukan oleh J. L. Fernández-Alemán et al. (2019) [8] yang menyatakan bahwa metode *ensemble learning* menghasilkan akurasi yang lebih baik daripada *individual learning*. Meskipun demikian, pernyataan ini membutuhkan agregasi bukti yang dilaporkan dalam literatur melalui tinjauan literatur yang sistematis. Begitu pun juga dalam penelitian yang dilakukan oleh P. S. Mung and S. Phyu (2020) [7] yang melakukan penggabungan hasil dari tiga algoritma *machine learning* untuk



lisis *big data* pelayanan kesehatan dan memperoleh hasil yang signifikan bahwa *ensemble learning* memiliki tingkat akurasi yang maksimal jika dibandingkan dengan metode *individual learning* atau *individual algorithm*.

Metode *ensemble learning* ini sendiri terbagi menjadi dua, yaitu *heterogeneous ensemble learning* dan *homogeneous ensemble learning* (Y. Feng et al. 2022) [9]. Dalam metode *heterogeneous ensemble learning*, kumpulan *base learner* terdiri dari beberapa jenis algoritma, dan strategi khusus digunakan untuk mengintegrasikan semua *base learner*. Kemudian untuk *homogeneous ensemble learning* sendiri mengacu pada fakta bahwa semua *base learner* dihasilkan menggunakan algoritma klasifikasi yang sama (Y. Feng et al. 2022) [9]. Kedua jenis metode *ensemble learning* tersebut sering digunakan dalam proses klasifikasi data atau regresi untuk melakukan prediksi, terkhususnya pada data dengan ukuran yang besar.

Alur pembentukan dari kedua metode *ensemble learning* ini pada umumnya sama. Tetapi khusus untuk *homogeneous ensemble learning*, karena konsepnya memadukan algoritma *machine learning* yang sama, maka menggunakan *training dataset* yang sama untuk setiap algoritma akan sia-sia karena akan menghasilkan keluaran model yang sama pula. Oleh karena itu untuk *training dataset* yang digunakan harus dilakukan proses pengambilan sampel secara acak (*Random Sampling*) terlebih dahulu menggunakan metode *Bagging (Bootstrap Aggregating)* (X. LI et al. 2019) [10] agar *training data* yang dimasukkan nantinya ke dalam setiap algoritma juga berbeda-beda, sehingga menghasilkan model terlatih yang berbeda-beda juga. *Bagging* adalah teknik *ensemble* yang sangat populer, teknik ini melakukan pengambilan sampel data dengan penggantian yang mengarah ke pemilihan titik data secara acak dan teknik ini juga dikenal karena kemampuannya untuk mengurangi *varians sampling*, yaitu *overfitting* (S. Alelyani 2021)(L. Li et al. 2019) [11][12].

Selain itu, pada proses akhir dari metode *ensemble learning* akan terdapat proses *voting* untuk mendapatkan kesimpulan prediksi atau final klasifikasi dari beberapa gabungan algoritma pengklasifikasi. Umumnya *ensemble learning* akan menggunakan metode *hard voting* dan *soft voting* dalam proses *voting* final klasifikasinya, tapi untuk penelitian ini sendiri akan berfokus pada penggunaan

weighted soft voting dalam melakukan proses *voting* kesimpulan final klasifikasi. Pada metode *simple soft voting*, akan menghasilkan nilai klasifikasi dengan merata-ratakan probabilitas yang dihitung dengan teknik individual



(A. Jenasamanta and S. Mohapatra 2022) [36], serta metode *simple soft voting* ini juga sudah terbukti lebih efektif daripada *hard voting / majority voting* dalam klasifikasi karena memprioritaskan *extremely confident votes* dan memasukkan signifikansi masing-masing pengklasifikasi ke dalam keputusan akhir (A. Taha 2021) [37]. Sedangkan untuk metode *weighted soft voting* yang digunakan pada penelitian ini akan mempertimbangkan nilai bobot setiap *individual classifier* dalam proses *voting* klasifikasi akhirnya. Dalam penentuan nilai bobot, digunakan algoritma pengoptimalan *Particle Swarm Optimization* (PSO) yang akan mencari kombinasi nilai bobot yang paling optimal untuk setiap *individual classifier*. Nilai bobot yang telah diperoleh tersebut kemudian akan digabungkan ke dalam metode *simple soft voting* untuk memperoleh hasil klasifikasi akhir.

Melihat kegunaan kedua jenis metode *ensemble learning* di atas yang sangat bermanfaat untuk proses klasifikasi dan ternyata lebih baik dibandingkan menggunakan *individual learning*, maka pada penelitian ini akan diusulkan penerapan metode *ensemble learning* dalam melakukan analisis sentimen. Metode *ensemble learning* ini akan digunakan untuk melakukan klasifikasi data *tweet* ke dalam sentimen positif, netral, atau negatif. Kemudian akan dilakukan juga analisis perbandingan untuk melihat performa dari kedua metode *ensemble learning* (*heterogeneous* dan *homogeneous ensemble learning*) dalam melakukan analisis sentimen pada data *tweet* publik terkait pembangunan Ibu Kota Negara Baru yang diambil dari sosial media Twitter. Sehingga nantinya dapat ditentukan jenis *ensemble learning* mana yang memiliki performa yang lebih baik dalam melakukan analisis sentimen.

1.2 Rumusan Masalah

Berdasarkan latar belakang masalah yang ada, maka rumusan masalah dari penelitian ini adalah:

1. Bagaimana meningkatkan performa dari algoritma *individual learning* dalam melakukan analisis sentimen menggunakan metode *ensemble learning*?



mana menentukan jenis *ensemble learning* yang memiliki performansi dalam proses analisis sentimen?

1.3 Tujuan Penelitian

Tujuan yang ingin dicapai pada penelitian ini adalah:

1. Menghasilkan model *ensemble learning* yang dapat meningkatkan performa dari algoritma *individual learning* dalam melakukan analisis sentimen *tweet* publik terkait pembangunan Ibu Kota Negara Baru Indonesia.
2. Memperoleh sebuah hasil analisis perbandingan antara *heterogeneous ensemble learning* dan *homogeneous ensemble learning* dari segi performa dalam melakukan analisis sentimen.

1.4 Manfaat Penelitian

Manfaat yang dapat diperoleh dari penelitian ini adalah sebagai berikut:

1. Mengetahui bagaimana sentimen atau opini masyarakat mengenai pembangunan Ibu Kota Negara Baru “Nusantara”.
2. Mengetahui jenis *ensemble learning* mana yang memiliki performa paling baik dalam melakukan proses analisis sentimen.
3. Bagi peneliti-peneliti selanjutnya atau mahasiswa yang memiliki penelitian yang relevan, dapat menjadi referensi dalam menentukan jenis *ensemble learning* apa yang sebaiknya mereka gunakan untuk proses analisis sentimen.
4. Memberikan pemahaman baru kepada peneliti mengenai analisis sentimen menggunakan metode *heterogeneous ensemble learning* dan *homogeneous ensemble learning*.

1.5 Batasan Masalah

Adapun batasan masalah dari penelitian ini adalah:

1. Data *tweet* yang akan diambil untuk dijadikan sebagai *dataset* adalah *tweet* dalam bahasa Indonesia.
 2. Algoritma individual yang akan digunakan dalam *ensemble learning* hanya 3, yaitu *Naïve Bayes*, *SVM*, dan *Logistic Regression*.
 3. Jenis sentimen hanya diklasifikasikan ke dalam tiga kelas, yaitu positif, netral, negatif.
- a sosial yang digunakan untuk mengambil data opini masyarakat mengenai pembangunan Ibu Kota Negara Baru adalah Twitter.



BAB II TINJAUAN PUSTAKA

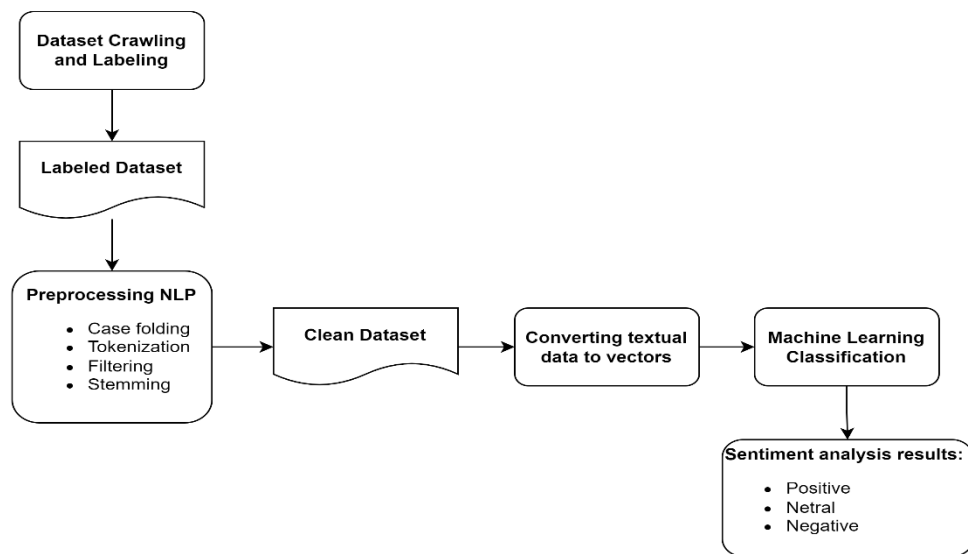
2.1 Kajian Pustaka

2.1.1 Analisis sentimen

Pada umumnya seseorang menginginkan pendapat dari orang lain sebagai masukan untuk menentukan keputusan. Pendapat ini dapat dilakukan dengan bertanya secara langsung, akan tetapi bertanya secara langsung membutuhkan lebih banyak waktu dan tenaga. Cara lain yang lebih mudah adalah dengan mendapatkan opini dari Twitter. Dengan memanfaatkan sosial media Twitter, kita dapat memperoleh opini banyak orang mengenai suatu hal atau suatu kategori melalui *tweet* mereka. Cara mendapatkan opini sesuai kategori yang diinginkan bisa menggunakan analisis sentimen (Kusrini and M. Mashuri 2019) [2]. Analisis Sentimen adalah metode untuk menentukan apakah sentimen dalam teks itu positif, netral, atau bahkan negatif. Ini juga dikenal sebagai polaritas material atau penambahan opini (L. Mandloi and R. Patel 2020) [3].

Analisis sentimen ini dapat membantu mendeteksi pendapat orang, sehingga memungkinkan pengambil keputusan untuk melacak perubahan dalam sentimen publik atau pelanggan mengenai entitas, aktivitas, produk, teknologi, dan layanan (M. T. H. K. Tusar and M. T. Islam 2021) [6]. Proses analisis sentimen akan melibatkan dua fase, fase pertama terkait dengan *Natural Language Processing* (NLP) dan fase kedua terkait dengan teknik *Machine Learning* (ML). Fase NLP adalah untuk mengetahui semantik kalimat, sedangkan ML digunakan untuk tujuan klasifikasi (A. Tareq and N. Hewahi 2021) [13]. Dalam melakukan analisis sentimen, terkhususnya melalui *tweet*, ada beberapa tahapan yang perlu dilalui, mulai dari pengumpulan dan pelabelan data *tweet*, melakukan tahapan pra-pemrosesan data, mengubah data tekstual ke dalam bentuk vektor, hingga proses klasifikasi sentimen menggunakan algoritma *machine learning*. Gambar 1 berikut merupakan gambaran umum dari tahapan analisis sentimen:





Gambar 1 Tahapan analisis sentimen

a. *Crawling dan labeling data*

Pada tahap awal akan dilakukan pengumpulan data *tweet* dari Twitter menggunakan Twitter API. Setelah data *tweet* dikumpulkan, data *tweet* akan diberikan label positif, netral, dan negatif sebagai jenis sentimen dari *tweet*-nya.

b. *Preprocessing*

Data yang diperoleh dari sosial media umumnya berantakan dan tidak terstruktur, sehingga perlunya dilakukan pembersihan terhadap data *tweet* yang telah diberi label tadi. Ada beberapa teknik yang dapat digunakan dalam tahapan pra-pemrosesan diantaranya adalah *case folding*, *tokenization*, *stopword removal/filtering*, dan *stemming* (D. A. Damaratih 2021) [5].

- 1) *Case Folding*, yaitu merubah setiap huruf pada kata menjadi huruf kecil atau *lower case*.
- 2) *Tokenization*, yaitu merupakan tahapan untuk memisah-misahkan kata atau membuat kata pada kalimat menjadi potongan-potongan per kata.
- 3) *Filtering*, yaitu proses untuk menghilangkan kata-kata yang kurang bermanfaat atau bermakna (*stopword*).
- 4) *Stemming*, yaitu proses untuk mengubah setiap kata menjadi kata dasar.



ubah data tekstual ke vektor

ata tekstual tidak dapat dipahami oleh model *machine learning*, sehingga perlu memasukkan data numerik agar model *machine learning* dapat pahamnya. Untuk melakukannya, kita harus mengubah data tekstual ke

dalam bentuk vektor menggunakan teknik *Natural Language Processing* untuk vektorisasi. Ada dua teknik *Natural Language Processing* yang populer untuk vektorisasi, yaitu *Bag-of-Words* (BoW) dan *Term Frequency-Inverse Document Frequency* (TF-IDF) (M. T. H. K. Tusar and M. T. Islam 2021) [6]. Untuk penelitian ini sendiri, akan digunakan metode TF-IDF.

d. Klasifikasi

Tahapan terakhir adalah proses pengklasifikasian *tweet* ke dalam sentimen positif, netral, dan negatif, dan proses ini akan dilakukan dengan menggunakan algoritma *machine learning*.

Telah banyak penelitian yang dilakukan terkait dengan analisis sentimen ini, berikut beberapa penelitian terkait dengan analisis sentimen:

A. Poornima and K. S. Priya (2020) [14] melakukan penelitian yang berfokus pada membandingkan kinerja algoritma pembelajaran mesin yang berbeda dalam melakukan analisis sentimen data Twitter. Analisis sentimen dari *tweet* ini membantu untuk menemukan apakah sentimen *tweet* pada produk tertentu, peristiwa, dll., positif atau negatif. Data yang digunakan dalam penelitian ini merupakan data *tweet* opini teks tentang produk atau acara, yang berisikan kelompok kata, emotikon, simbol, URL, dan referensi ke orang-orang. Metode atau algoritma klasifikasi pembelajaran mesin yang diusulkan dan akan dibandingkan adalah *Logistic regression*, SVM, dan *Multinomial Naïve Bayes*. Hasilnya menunjukkan bahwa tingkat akurasi untuk algoritma *Logistic regression* merupakan yang tertinggi yaitu 86.23%, kemudian diikuti oleh SVM sebesar 85.69%, dan yang paling rendah adalah *Multinomial Naïve Bayes* sebesar 83.54%.

M. T. H. K. Tusar and M. T. Islam (2021) [6] melakukan penelitian yang membahas tentang proses membandingkan algoritma *Machine Learning* dan teknik NLP yang diterapkan untuk menemukan pendekatan terbaik dalam melakukan analisis sentimen pada *dataset* Twitter *US Airline Sentiment* yang memiliki sekitar 14640 *record* dan 15 atribut. Metode yang digunakan dalam penelitian ini adalah *Bag-of-Words* (BoW) dan TF/IDF untuk teknik NLP, serta *Support Vector*

Multinomial Naive Bayes, *Random Forest*, dan *Logistic Regression* untuk *machine learning*-nya. Dalam prosesnya, akan dibandingkan tingkat presisi, *recall*, dan F1-Score dari 2 teknik NLP yaitu *Bag-of-Words* (BoW)



dan TF/IDF yang digabungkan dengan setiap algoritma *machine learning*. Hasil yang diperoleh adalah ketika teknik *Bag-of-Words* (BoW) digabungkan dengan setiap algoritma *machine learning*, maka algoritma SVM dan *Logistic Regression* yang memiliki tingkat akurasi tertinggi, yaitu masing-masing mencapai 77%. Sama halnya ketika teknik TF/IDF digabungkan dengan setiap algoritma *machine learning*, yang diperoleh juga sama bahwa algoritma SVM dan *Logistic Regression* yang memiliki tingkat akurasi tertinggi, yaitu masing-masing mencapai 77%.

Kusrini and M. Mashuri (2019) [2] melakukan penelitian yang membahas tentang penggunaan metode *Lexicon* dan *Polarity Multiplication* dalam melakukan proses analisis sentimen pada kumpulan data secara umum yang diambil dari Twitter menggunakan Twitter API. Dalam proses analisis sentimen ini dilakukan beberapa teknik NLP dalam tahapan pra-pemrosesan teksnya, teks yang sudah bersih selanjutnya ditentukan jenis sentimennya menggunakan aturan perkalian bilangan bulat (*Polarity Multiplication*) dan pencocokan terhadap *Lexicon* yang disediakan. Pada penelitian ini hasil dari metode yang diusulkan dibandingkan dengan hasil analisis sentimen ketika menggunakan algoritma *machine learning*, dan diperoleh bahwa tingkat akurasi ketika menggunakan metode *Lexicon* dan *Polarity Multiplication* dalam melakukan analisis sentimen masih lebih rendah dibandingkan ketika menggunakan *machine learning*. Tingkat akurasi rata-rata ketika menggunakan metode yang diusulkan adalah 68%, sedangkan ketika menggunakan beberapa algoritma *machine learning* seperti *Naïve Bayes* dan SVM, diperoleh tingkat akurasi rata-rata untuk setiap algoritma masing-masing mencapai 86% dan 77%. Hal ini dikarenakan jumlah kata sifat dalam *Lexicon* masih belum lengkap. Selain itu, metode *Lexicon* juga lebih sederhana daripada metode yang tersedia dalam pembelajaran mesin.

L. Mandloi and R. Patel (2020) [3] melakukan penelitian yang membahas tentang bagaimana proses analisis sentimen dilakukan pada media sosial Twitter untuk memperoleh sentimen dari setiap *tweet* masyarakat menggunakan tiga algoritma *machine learning* dan membandingkan ketiga algoritma berdasarkan

akurasi dan presisi untuk melihat algoritma mana yang memberikan hasil. Dalam pengumpulan *dataset*-nya digunakan Twitter API. Adapun tiga yang digunakan dan dibandingkan dalam penelitian ini adalah *Naïve Bayes*,



SVM, dan *Maximum Entropy*. Hasil yang diperoleh adalah algoritma *Naïve Bayes* mendapatkan skor akurasi dan presisi tertinggi dibandingkan dua algoritma lain, yaitu 86% untuk akurasi dan 88.7% untuk presisi. Sedangkan untuk *Maximum Entropy* memperoleh skor akurasi 82.6% dan presisi 84.05%, dan *SVM* memiliki skor akurasi dan presisi terendah, yaitu 74.6% untuk akurasi dan 75.9% untuk presisi.

S. A. El Rahman et al. (2019) [15] melakukan sebuah penelitian untuk menyajikan model yang dapat melakukan analisis sentimen dari data nyata yang dikumpulkan dari Twitter. Data di Twitter sangat tidak terstruktur sehingga sulit untuk dianalisis, sehingga diusulkan model gabungan dari algoritma *supervised* dan *unsupervised machine learning*. Pengumpulan *dataset* dilakukan menggunakan Twitter API. Data yang dikumpulkan merupakan teks berbahasa Inggris dan dikumpulkan pada dua subjek yaitu McDonalds dan KFC untuk menunjukkan restoran mana yang lebih populer berdasarkan ulasan baik/buruk dari masyarakat. Algoritma *unsupervised learning* dengan model berbasis *lexicon* digunakan untuk melabeli data *tweet* secara otomatis. Dataset yang sudah diberi label selanjutnya dilatih menggunakan algoritma *supervised learning*, seperti *Naïve Bayes*, *SVM*, *maximum entropy*, *decision tree*, *random forest* dan *bagging*. Berdasarkan pengujian *Cross Validation*, diperoleh hasil yang menunjukkan bahwa *maximum entropy* adalah model klasifikasi terbaik untuk analisis sentimen menggunakan data KFC dan McDonalds jika dibandingkan dengan model lainnya. Dimana untuk data McDonalds, model ini memperoleh skor akurasi *cross validate* 74% dan untuk data KFC, model ini memperoleh skor akurasi *cross validate* 78%.

2.1.2 Ensemble learning

Metode *Ensemble learning* adalah salah satu pendekatan yang mencoba untuk memperkuat tugas pembelajaran mesin, dan beberapa penelitian telah menunjukkan bahwa metode ini biasanya memiliki kinerja yang lebih tinggi dibanding dengan *individual learning* (P. Puvar et al. 2021) [21]. Metode *ensemble learning* ini merupakan sebuah metode yang menggabungkan beberapa algoritma *machine learning* untuk mendapatkan performa yang lebih baik, yang dimana ide dasar dari *learning* ini adalah menerapkan model pembelajaran yang berbeda-beda untuk mendapatkan klasifikasi atau hasil yang lebih baik dari *individual learning*.



Berikut beberapa penelitian yang telah dilakukan menggunakan metode *ensemble learning* dan membuktikan bahwa metode ini dapat memiliki kinerja lebih baik dari *individual learning*:

P. S. Mung and S. Phyu (2020) [7] melakukan penelitian yang membahas tentang proses analisis yang akan dilakukan terhadap *big data* layanan kesehatan. Untuk meningkatkan akurasi dari metode *individual learning* yang telah banyak digunakan untuk analisis data pelayanan kesehatan, diusulkan pendekatan pembelajaran mesin yang disebut *ensemble learning*, di mana hasil dari tiga algoritma pembelajaran mesin digabungkan yaitu *Naïve Bayesian*, *Decision Tree* dan *K-NN*. Untuk menggabungkan hasil akurasi digunakan metode *soft voting*. Pengujian dilakukan dengan membandingkan hasil akurasi *ensemble learning* dengan hasil akurasi *individual learning* setiap algoritma, yang di uji pada 4 dataset berbeda. Hasilnya, dari keempat dataset yang diuji, metode *ensemble learning* memperoleh tingkat akurasi sebesar 99.93%, 93.56%, 98.06% dan 98.82%, yang dimana tingkat akurasi ini lebih tinggi dibandingkan dengan *individual learning* setiap algoritma. Sehingga dapat disimpulkan bahwa tingkat akurasi yang diperoleh akan lebih baik jika setiap algoritma digabungkan menggunakan metode *Ensemble learning* dibandingkan dengan *individual learning*.

S. Mehta et al. (2019) [22] melakukan penelitian yang membahas tentang pemanfaatan metode *Ensemble Learning* dalam memprediksi pasar saham untuk memutuskan estimasi masa depan saham organisasi atau instrumen terkait uang dan lainnya yang dipertukarkan dalam perdagangan moneter, sehingga dapat menghasilkan keuntungan investasi yang tinggi. Model yang dibangun dalam penelitian ini adalah *Weighted ensemble model*, yang merupakan gabungan antara metode *weighted support vector regression* (SVR), *Long-short term memory* (LSTM), dan *Multiple Regression*. Model *ensemble learning* tersebut akan dibandingkan dengan penggunaan model secara individu. Hasilnya menunjukkan tingkat akurasi model *ensemble* lebih tinggi yaitu 99.12% dibandingkan dengan penggunaan model secara individu.



Sanabila and W. Jatmiko (2018) [23] melakukan penelitian yang berfokus pada evaluasi kinerja *ensemble learning* dalam menangani *imbalanced data*. *Imbalanced data* merupakan masalah khusus dalam tugas klasifikasi dimana

distribusi kelas tidak seragam. *Resampling* (SMOTE dan ENN) digunakan untuk meningkatkan kinerja *classifier*. Empat metrik diterapkan untuk evaluasi kinerja yaitu *precision*, *recall*, *specificity* dan *F-1 score*. Penelitian ini menggunakan tiga metode *Ensemble* yaitu *Boosting*, *Random Forest*, dan *Bagging*. Sedangkan *Naïve Bayes* dan *Logistic Regression* digunakan sebagai metode *baseline*. Penelitian ini dilakukan dalam tiga rasio data *training-testing* yaitu 70-30, 80-20, dan 90-10. Hasilnya *Bagging* memiliki kinerja yang unggul dalam *imbalanced dataset* dan *resampling* dibandingkan dengan *baseline classifiers* (*Naïve Bayes* dan *Log Regression*) dan *ensemble learning* lainnya (*Boosting* dan *Random Forest*). Selain itu, kombinasi SMOTE dan ENN berhasil meningkatkan kinerja klasifikasi dan menghindari bias pada kelas mayoritas

J. Tie, X. Lei and Y. Pan (2022) [24] melakukan penelitian yang berfokus pada pemanfaatan *machine learning* dalam pengidentifikasian hubungan antara metabolit dan penyakit untuk membantu memahami patogenesis penyakit yang sangat penting dalam mendiagnosis dan mengobati penyakit. Penelitian ini mengusulkan algoritma prediksi asosiasi penyakit metabolit baru berdasarkan gabungan antara *DeepWalk* dan *random forest* (DWRF). *DeepWalk* digunakan untuk mengekstrak fitur metabolit berdasarkan jaringan asosiasi gen metabolit, kemudian algoritma *random forest* digunakan untuk menyimpulkan asosiasi penyakit-metabolit. Hasil eksperimen menunjukkan bahwa metode DWRF memiliki performa yang lebih baik dibandingkan algoritma klasifikasi individual lainnya seperti *logistic regression*, SVM, *Gaussian naive bayes* (GNB), KNN dan *AdaBoost*, dimana DWRF mendapatkan nilai *area under the precise-recall curve* (AUPR) yang mencapai 97%, nilai *F1-Score* 90%, nilai akurasi 89%, nilai *specificity* 90%, nilai *recall* 89%, dan nilai presisi 90%.

J. L. Fernández-Alemán et al. (2019) [8] melakukan penelitian yang mengeksplorasi penggunaan metode klasifikasi *ensemble* dalam konteks penyakit diabetes. Deteksi dini diabetes yang akurat merupakan primordial untuk menghindari komplikasi kesehatan dengan memberikan terapi yang tepat kepada



Untuk mengatasi kelemahan dari sistem *computer-aided* /detection (CAD), proses *ensemble learning* diusulkan. *Bibliographical* s *IEEE Xplore*, Scopus, *ACM Digital Library* dan PubMed digunakan

untuk memilih makalah yang melaporkan metode *ensemble classification* pada penyakit diabetes. Sebanyak 107 makalah akhirnya dipilih setelah melalui proses seleksi studi. Metode ensemble diterapkan untuk diabetes pada tahun 2003 untuk pertama kalinya. *Homogeneous ensembles* adalah yang paling umum digunakan dalam literatur. Selain itu, *decision tree* dan SVM adalah teknik yang paling banyak digunakan untuk membangun *homogeneous* dan *heterogeneous ensembles*. Metode penggabungan yang paling sering ditemukan adalah *majority voting rule*. Temuan pada *review* yang dilakukan pada penelitian ini menunjukkan bahwa metode *ensemble classification* menghasilkan akurasi yang lebih baik daripada *single classifiers*.

2.1.3 Naïve bayes

Naïve Bayes adalah algoritma *supervised machine learning* yang menggunakan teorema Bayes untuk masalah klasifikasi dan merupakan salah satu algoritma yang paling sederhana dan paling kuat untuk membantu menciptakan model pembelajaran mesin yang cepat dalam membuat prediksi. Algoritma ini sebagian besar digunakan untuk permasalahan klasifikasi teks, seperti analisis teks, analisis sentimen, mengklasifikasikan artikel, dan lain sebagainya (L. Mandloi and R. Patel 2020) [3].

$$P(A|B) = \frac{P(B|A) * P(A)}{P(B)}$$

Dimana $P(A|B)$ biasa disebut *posterior probability*, yaitu probabilitas terjadinya kejadian A ketika terjadi kejadian B; $P(B|A)$ biasa disebut *likelihood probability*, yaitu probabilitas terjadinya kejadian B ketika terjadi kejadian A; $P(A)$ biasa disebut *prior*, yaitu probabilitas terjadinya kejadian A; dan $P(B)$ biasa disebut *evidence*, yaitu probabilitas terjadinya kejadian B.

2.1.4 SVM (Super Vector Machine)

SVM dapat digunakan untuk mengklasifikasikan *tweet* sebagai positif dan negatif dari vektor *input* dan kemudian menghitung keakuratan sentimen dengan bantuan *term frequency* (A. Poornima and K. S. Priya 2020) [14]. SVM bekerja

mengklasifikasikan data dalam kelas yang berbeda dengan menggambar garis pemisah antara kelas-kelas tersebut, yang membedakan antara kelas-kelas yang berbeda yang diplot dalam ruang n-dimensi (L. Mandloi and R. Patel 2020) [3].



2.1.5 Logistic regression

Logistic Regression adalah algoritma analisis prediktif yang didasarkan pada konsep probabilitas. *Logistic Regression* menggunakan *cost function* yang lebih kompleks yang disebut ‘*Sigmoid function*’. Hipotesis *Logistic Regression* cenderung membatasi *cost function* antara 0 dan 1 (A. Poornima and K. S. Priya 2020) [14].

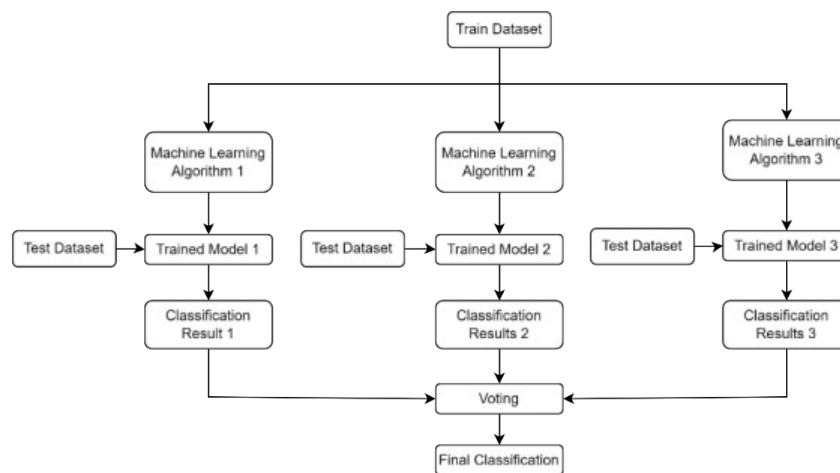
$$\text{Sigmoid Function } f(x) = \frac{1}{1 + e^{-x}}$$

2.1.6 Heterogeneous dan homogeneous ensemble learning

Metode *ensemble learning* dapat dibagi menjadi dua, yaitu *heterogeneous ensemble learning* dan *homogeneous ensemble learning* (Y. Feng et al. 2022) [9].

a. Heterogeneous ensemble learning

Heterogeneous ensemble learning merupakan pepaduan sejumlah model *machine learning* yang berbeda untuk mendapatkan sebuah hasil yang lebih baik dari proses klasifikasi atau regresi. Secara sederhana metode ini terdiri dari anggota yang memiliki *base learner* atau algoritma yang berbeda (P. Puvar et al. 2021) [21]. Berikut adalah gambaran umum dari tahapan *heterogeneous ensemble learning*:



Gambar 2 Tahapan *heterogeneous ensemble learning*

Pada Gambar 2 di atas memperlihatkan bagan tahapan proses *heterogeneous ensemble learning* yang mana dapat dilihat bahwa *train dataset* dimasukkan ke beberapa algoritma *machine learning* yang berbeda, dan masing-masing akan menghasilkan model yang telah terlatih. *Test dataset* yang dimasukkan ke dalam model, masing-masing akan menghasilkan hasil klasifikasi dan metode digunakan untuk memperoleh hasil klasifikasi final.



Berikut adalah beberapa penelitian yang telah dilakukan terkait dengan metode *heterogeneous ensemble learning*:

C. Zhao et al. (2020) [25] melakukan penelitian yang mengusulkan *heterogeneous stacking-based ensemble learning framework* untuk memperbaiki dampak ketidakseimbangan kelas pada deteksi *spam* di jejaring sosial, melihat banyaknya perilaku jahat *spammer* jejaring sosial yang telah mengancam keamanan informasi pengguna biasa. *Framework* yang diusulkan terdiri dari dua komponen utama, yaitu *base module* dan *combining module*. Untuk membentuk *base module* dari *framework*, digunakan *base classifiers*, yaitu SVM, CART, *Naïve Bayes*, KNN, *Random Forest*, dan *Linear Regression*. Hasil prediksi dari *base classifiers* ini digunakan sebagai *input* dari *combining module*. Dalam *combining module*, diterapkan *cost-sensitive learning* ke dalam pelatihan *deep neural network* (DNN). Penelitian ini menggunakan *dataset* yang berisi 600 juta *tweet* di mana 6,5 juta di antaranya adalah *malicious tweets*. Hasil eksperimen menunjukkan bahwa *framework* yang diusulkan secara efektif meningkatkan tingkat deteksi *spam* pada kumpulan data yang tidak seimbang.

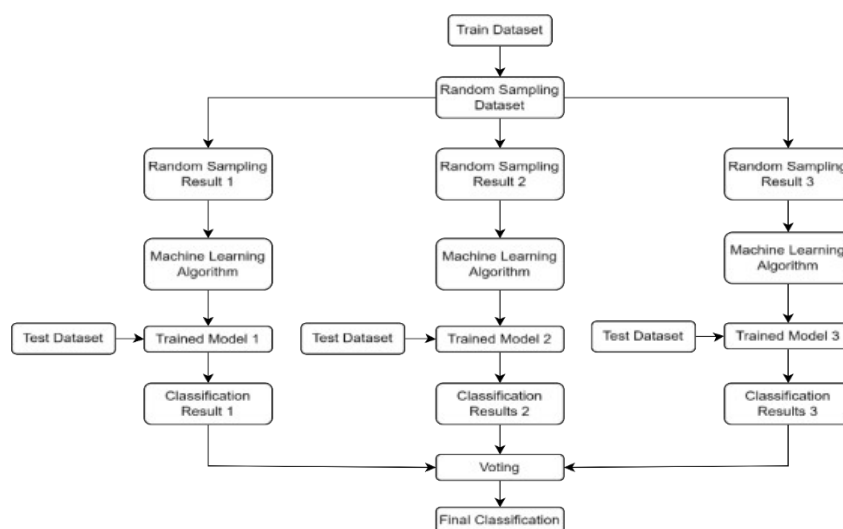
Y. Feng et al. (2022) [9] melakukan penelitian yang membahas tentang bagaimana melakukan prediksi kelangsungan hidup yang akurat dari pasien penderita neuroblastoma, sehingga diusulkan metode *heterogeneous ensemble learning* untuk memprediksi kelangsungan hidup pasien neuroblastoma dan mengekstrak aturan keputusan dari metode yang diusulkan untuk membantu dokter dalam membuat keputusan dan rencana perawatan untuk penyakit ini. Dalam membangun metode *heterogeneous* ini, dikembangkan lima *heterogeneous base learner* yang terdiri dari *decision tree*, *random forest*, *support vector machine based on genetic algorithm*, *extreme gradient boosting*, dan *light gradient boosting machine*. *Dataset* yang digunakan dalam penelitian ini merupakan *dataset* neuroblastoma dari *Therapeutically Applicable Research To Generate Effective Treatments* (TARGET), yang berfokus pada tumor anak-anak, dan dikelola oleh Program Evaluasi Terapi Kanker dan Genomik Kanker NCI. Hasil eksperimen



akan bahwa metode yang diusulkan mencapai akurasi 91.64%, *recall* dan AUC 91.35%, serta secara signifikan lebih baik daripada metode *machine learning*.

b. *Homogeneous ensemble learning*

Homogeneous ensemble learning merupakan pemaduan sejumlah model *machine learning* yang sama atau sejenis untuk mendapatkan sebuah hasil yang lebih baik dari proses klasifikasi atau regresi. Secara sederhana metode ini terdiri dari anggota yang hanya memiliki satu *base learner* atau algoritma (P. Puvar et al. 2021) [21]. Khusus untuk metode *homogeneous* ini, karena menggunakan anggota algoritma yang sama, maka *training dataset* yang akan dilatih harus melalui proses *random sampling* terlebih dahulu agar *training data* yang dimasukkan nantinya ke dalam algoritma juga berbeda-beda, sehingga model terlatih yang dihasilkan akan berbeda-beda pula. Berikut adalah gambaran umum dari tahapan *homogeneous ensemble learning*:



Gambar 3 Tahapan *homogeneous ensemble learning*

Pada Gambar 3 di atas memperlihatkan bagan tahapan proses *homogeneous ensemble learning* yang mana dapat dilihat bahwa *train dataset* terlebih dahulu harus melalui proses *random sampling*, salah satu metode *random sampling* yang populer adalah metode *Bagging* (X. LI et al. 2019) [10]. Masing-masing *dataset* yang diperoleh dari hasil proses *random sampling*, kemudian dimasukkan ke dalam sebuah algoritma *machine learning*, untuk mendapatkan model terlatih. *Test dataset* yang dimasukkan ke dalam setiap model terlatih, masing-masing akan menghasilkan hasil klasifikasi dan metode *voting* digunakan untuk memperoleh klasifikasi final.

Ada beberapa penelitian yang telah dilakukan terkait dengan metode *homogeneous ensemble learning*:



Murni et al. (2019) [26] melakukan penelitian analisis sentimen terhadap ulasan *TripAdvisor* yang merupakan *web* ulasan wisata terkemuka. Hal ini dilakukan karena ditemukan adanya ketidakkonsistenan antara komentar wisatawan terhadap objek wisata dengan ulasan atau penilaian mereka yang digambarkan dengan peringkat bintang. Dalam melakukan analisis sentimen, penelitian ini mengusulkan metode *hybrid*, yaitu metode berbasis *lexical* dan metode berbasis *machine learning*. Dalam metode berbasis *lexical* digunakan *SenticNet* untuk menentukan label perwakilan setiap ulasan. Adapun untuk metode berbasis *machine learning* diterapkan *homogeneous ensemble classifiers*, yaitu *bagged decision trees*, *logistic model tree*, dan *bagged multi-layer perceptron algorithms* untuk menentukan sentimen dari data review. Hasil eksperimen review *TripAdvisor* terhadap 55 objek wisata di Bali menunjukkan bahwa *Homogeneous Ensemble Random Forest classifier* memberikan akurasi rata-rata tertinggi, yaitu 98.1267%.

K. S. Gan et al. (2018) [27] melakukan penelitian yang membahas tentang bagaimana meramal harga penutupan pasar saham CIMB menggunakan *single* dan *multiple homogeneous ANN*. Kinerja dari *feedforward neural network* (FFNN) dan *ensemble FFNN* (ENN) dalam melakukan peramalan harga penutupan saham CIMB dibandingkan dan dianalisis. Data yang dikumpulkan untuk penelitian ini meliputi informasi saham (harga pembukaan, harga penutupan, harga tertinggi, harga terendah dan volume perdagangan) dan informasi yang mencerminkan pasar seperti suku bunga, KLCI dan nilai tukar mata uang (USD, EUR dan SGD). Kumpulan data untuk semua parameter diambil dari Januari 2000 hingga Juni 2015. Hasil penelitian menunjukkan bahwa baik FFNN dan ENN berkinerja baik prediksi, dengan mencapai akurasi prediksi lebih dari 90%. Namun dalam semua kriteria evaluasi kinerja ENN lebih baik dari FFNN, terkhususnya pada kriteria evaluasi persentase akurasi prediksi 100%, dimana ENN memperoleh persentase sebesar 59.37% dan FFNN hanya 58.1%.

2.1.7 Weighted soft voting

Seperti yang diketahui, Konsep di balik metode *ensemble learning* adalah untuk meningkatkan kinerja model prediktif dengan menggabungkan prediksi dari pengklasifikasi (K. Kurnianingsih et al. 2020) [18]. Untuk mendapatkan prediksi dari beberapa gabungan algoritma pengklasifikasi, dapat



dilakukan dengan metode *voting*. Metode *voting* salah satu cara yang efisien untuk menggabungkan prediksi dari beberapa algoritma pembelajaran mesin (A. Jenasamanta and S. Mohapatra 2022) [36]. Salah satu metode *voting* yang banyak digunakan dan terbukti efektif dalam klasifikasi *ensemble* adalah *soft voting*.

Pada metode *simple soft voting*, prediksi akhir adalah kelas dengan probabilitas tertinggi yang dirata-ratakan atas semua *individual classifier* (A. Jenasamanta and S. Mohapatra 2022) [36]. Penelitian yang dilakukan oleh A. Taha (2021) [37] juga telah membuktikan bahwa metode *simple soft voting* lebih efektif daripada *hard voting* / *majority voting* dalam klasifikasi karena memprioritaskan *extremely confident votes* dan memasukkan signifikansi masing-masing pengklasifikasi ke dalam keputusan akhir. Secara sederhana, rumus *soft voting* dapat ditulis sebagai berikut:

$$P_{final}(K_i) = \frac{1}{m} \sum_{j=1}^m P_{M_j}(K_i)$$

Ket:

$P_{final}(K_i)$ = probabilitas akhir untuk kelas K_i berdasarkan *soft voting*.

$P_{M_j}(K_i)$ = probabilitas prediksi kelas K_i dari model M_j .

m = total model.

Khusus untuk metode *weighted soft voting* sendiri, juga akan menggunakan *simple soft voting*, tetapi dengan penambahan bobot (*weight*) dari setiap *individual classifier* ke dalam proses perhitungan *soft voting*-nya. Setiap *individual classifier* akan memiliki bobot yang berbeda-beda, tergantung performanya dalam melakukan proses klasifikasi (R. Wardoyo et al. 2020) [38]. Dengan kata lain, *weighted soft voting* akan mempertimbangkan bobot setiap *individual classifier* dalam proses *soft voting* yang dilakukan. Rumus dari *weighted soft voting* sendiri dapat ditulis sebagai berikut:

$$P_{final}(K_i) = \frac{1}{m} \sum_{j=1}^m w_j \cdot P_{M_j}(K_i)$$

Ket:

$P_{final}(K_i)$ = probabilitas akhir untuk kelas K_i berdasarkan *soft voting*.

$P_{M_j}(K_i)$ = probabilitas prediksi kelas K_i dari model M_j .

w_j = bobot untuk model M_j .

m = total model.



Berikut beberapa penelitian yang telah dilakukan terkait *ensemble learning* yang menggunakan metode pembobotan (*weighted*) dalam proses *voting*-nya dan membuktikan bahwa dengan mempertimbangkan bobot setiap *individual classifier* dalam melakukan *voting* akan memberikan hasil yang baik untuk *final classification*-nya:

A. Taha (2021) [37] melakukan penelitian yang membahas tentang *website phishing* sebagai ancaman *cybersecurity*, dengan membuat halaman *website* untuk meniru halaman *website* asli untuk menipu orang dan mencuri informasi pribadi mereka, sehingga perlunya pendeteksian akurat *website phishing*. Penelitian ini mengusulkan pendekatan *intelligent ensemble learning* untuk deteksi *website phishing* berdasarkan *weighted soft voting* untuk meningkatkan deteksi *website phishing*. Lima algoritma *base classifier* digunakan dalam *ensemble learning*, yaitu *Decision Tree*, *Naïve Bayes*, *Random Forest*, *Gradient Boosting*, dan *Logistic Regression*, serta metode *weighted soft voting* berdasarkan statistik Kappa juga digunakan. Eksperimen dilakukan menggunakan kumpulan data *website phishing* yang tersedia untuk umum dari *UCI Machine Learning Repository*, yang terdiri dari 4898 *website phishing* dan 6157 *website* yang sah. Hasil yang diperoleh menunjukkan bahwa pendekatan pendekatan *intelligent ensemble learning* yang diusulkan untuk deteksi *website phishing* memperoleh tingkat akurasi tertinggi 95% dan *Area Under the Curve* (AUC) sebesar 98,8%.

R. Wardoyo et al. (2020) [38] melakukan penelitian yang membahas tentang usulan perancangan metode untuk memberi bobot/*weight* sebuah *single classifier* dengan berfokus pada penilaian kinerja yang dihitung berdasarkan hasil prediksi setiap kelas. Pembobotan dilakukan dengan menganalisis performa setiap *base classifier* yang diukur dari parameter akurasi, *recall*, presisi, dan *F-Measure*. *Dataset* yang digunakan merupakan dataset publik dan data *umbilical cord* yang diperoleh *UCI repository*, *Keel repository*, dan *umbilical cord image* pada penelitian sebelumnya. Metode *ensemble learning* dengan *weighted majority voting* digunakan dalam penelitian ini dan melibatkan lima *base classifier* untuk model

-nya, yaitu *Random Forest*, *Decision Tree* (C.45), *Gradient Boosting* *XGBoost*, dan *Bagging*. Nilai bobot untuk setiap model selalu berkorelasi inerja *single classifier* dimana kemampuan prediksi yang lebih baik akan



menghasilkan bobot yang lebih tinggi pula. Hasil dari penelitian menunjukkan bahwa metode yang diusulkan mampu meningkatkan akurasi dari *base classifier* dari penelitian sebelumnya. Selain itu, metode yang diusulkan juga meningkatkan kinerja *dataset umbilical cord* dengan akurasi rata-rata 86.1%, presisi 86%, *recall* 86%, dan *F-Measure* 86%.

A. Dogan and D. Birant (2019) [39] melakukan penelitian yang mengusulkan sebuah pendekatan baru yaitu *Weighted Majority Voting Ensemble* (WMVE) untuk mengevaluasi kinerja individu dari setiap *classifier* dalam *ensemble learning* dan menyesuaikan kontribusinya terhadap keputusan kelas. Dalam model *ensemble* yang diusulkan, setiap *classifier* akan memiliki tingkat pengaruh yang berbeda-beda terhadap prediksi akhir tergantung nilai bobotnya. Nilai bobot setiap *classifier* akan disesuaikan, jika *classifier* berhasil memprediksi dengan benar maka bobotnya akan semakin bertambah. Keputusan akhir dibuat dengan menjumlahkan semua *weighted voting* dan memilih kelas dengan agregat tertinggi. Algoritma klasifikasi yang digunakan dalam model *ensemble* yang diusulkan adalah C4.5 *decision tree*, SVM, KNN, *k-star*, dan *Naïve Bayes*. Hasil dari metode WMVE yang diusulkan dibandingkan dengan pendekatan *Simple Majority Voting Ensemble* (SMVE), serta dengan algoritma klasifikasi standar individu dalam hal akurasi klasifikasi. *Dataset* yang digunakan terdiri dari 28 *dataset benchmark* dari berbagai bidang seperti kedokteran, otomotif, keuangan dan sebagainya. Berdasarkan hasil akurasi rata-rata dari pengujian pada 28 *dataset*, pendekatan WMVE yang diusulkan memiliki skor akurasi terbaik dengan 90,17%. Hasil tersebut mengungguli pendekatan lain dalam 20 dari 28 *dataset*.

R. Islam et al. (2023) [40] melakukan penelitian yang membahas tentang bagaimana membangun sebuah sistem klasifikasi *malware* pada android dengan mengidentifikasi kategorinya menggunakan algoritma *Machine Learning*. *Dataset* yang digunakan dalam penelitian ini adalah CCCS-CIC-AndMal-2020, merupakan kumpulan data yang tersedia untuk umum yang diproduksi pada tahun 2020 oleh *Canadian Centre for Cyber Security* dan *Canadian Institute for Cybersecurity*, yang



iyak koleksi aplikasi Android dan sampel *malware*. Pendekatan *ensemble learning* dan *weighted voting* digunakan dalam penelitian ini dengan ungkan beberapa algoritma individual machine learning yang berbeda

untuk mode *ensemble*-nya, yaitu *Random Forest*, *K-NN*, *Multi-Level Perceptrons*, *Decision Trees*, *SVM*, dan *Logistic Regression*. Hasil metode *ensemble machine learning* dan *weighted voting* yang diusulkan pada penelitian ini kemudian dibandingkan dengan masing-masing individual algoritma *machine learning* tersebut. Hasilnya, metode yang diusulkan memperoleh akurasi sebesar 95% dan lebih unggul dibandingkan semua algoritma individual.

2.1.8 *Bagging (Bootstrap Aggregating)*

Bagging adalah pengklasifikasi yang digunakan untuk mengambil beberapa sampel acak dan menggunakan setiap sampel secara terpisah untuk membangun model prediksi (S. A. El Rahman et al. 2019) [15]. *Bagging* menggunakan *sampling by replacement* untuk membangun model dan mengeksploitasi banyak *independent learner* serta menggabungkannya menggunakan teknik *averaging* (H. R. Sanabila and W. Jatmiko 2018) [23]. Teknik *bagging* ini sendiri memiliki kemampuan untuk mengurangi *varians sampling* seperti *overfitting* dan mendapatkan solusi yang lebih stabil (S. Alelyani 2021)(L. Li et al. 2019) [11][12].

Berikut adalah beberapa penelitian yang telah dilakukan dengan memanfaatkan metode *Bagging*:

S. Alelyani (2021) [11] melakukan penelitian yang membahas tentang pendekatan *ensemble* berdasarkan teknik *bagging* untuk meningkatkan stabilitas pemilihan fitur dalam kumpulan data medis melalui pengurangan *varians* data. Percobaan dilakukan menggunakan empat set data *microarray* yang masing-masing memiliki dimensi tinggi dan ukuran sampel yang relatif kecil. Kemudian pada penelitian ini lima algoritma pemilihan fitur terkenal yang berbeda digunakan. Selain itu, dua algoritma klasifikasi juga digabungkan, yaitu KNN dan SVM. Dalam hampir semua kasus, metode yang diusulkan ditemukan untuk mengatasi pendekatan dasar dengan peningkatan stabilitas yang cukup besar mulai dari 20% hingga 50%. Selain itu, akurasi klasifikasi ditingkatkan di lebih dari setengah percobaan dan dipertahankan di sisanya.

L. Li et al. (2019) [12] melakukan penelitian yang berfokus pada pengembangan



embelajaran mesin yang sangat akurat dan adaptif untuk memprediksi isi NO_2 dan NO_x resolusi tinggi di wilayah geografis yang luas akan pengukuran dari berbagai sumber yang berisi sampel dengan

distribusi spatiotemporal dan pola pengelompokan yang heterogen, karena hal ini sangat penting untuk meningkatkan evaluasi efek kesehatan. Metode yang digunakan dalam penelitian ini adalah metode *comprehensive Kruskal-K-means* untuk mengelompokkan sampel pengukuran dari berbagai sumber yang heterogen. *Spatiotemporal cluster-based bootstrap aggregating (bagging)* dari model efek campuran dasar kemudian diterapkan. Selain itu digunakan juga teknik pembelajaran mesin *grid search* untuk menemukan interaksi optimal dari fungsi basis temporal dan skala efek spasial, yang bersama dengan kovariat spatiotemporal, secara memadai menangkap variabilitas spatiotemporal dalam NO_2 dan NO_x di tingkat negara bagian dan lokal. Hasil yang didapatkan adalah, dengan *cluster-based bagging* dari model dasar diperoleh prediksi yang kuat dengan *ensemble cross validation R2* sebesar 0.88 untuk NO_2 dan NO_x , RMSE sebesar 3.62-9.63ppb, dan RMSEIQR sebesar 0.28-0.37. Dalam pengujian *random sampling independent*, model mencapai kinerja yang sama kuatnya, yaitu R^2 sebesar 0.87-0.90, RMSE sebesar 3.97-9.69 ppb, dan RMSEIQR sebesar 0.21-0.27.

2.1.9 TF-IDF (Term Frequency - Inverse Document Frequency)

Metode TF-IDF menggunakan 2 parameter pembobotan yaitu pembobotan lokal $tf_{i,d}$. Ini adalah bobot yang diperoleh dari frekuensi kata “ t ” dalam dokumen “ d ” dan bobot global dengan menggunakan idf_t adalah bobot yang diperoleh dengan mempertimbangkan jumlah kata “ t ” (df_t) di semua dokumen N . Setelah itu, nilai dari bobot lokal dikalikan dengan nilai dari bobot global (F. A. Wenando et al. 2020) [4]. TF-IDF ini berfungsi untuk mencari istilah atau kata penting yang muncul dalam dokumen atau tweet berdasarkan frekuensinya (M. T. H. K. Tusar and M. T. Islam 2021) [6], atau secara sederhana metode ini membuat kata yang sering muncul memiliki nilai yang cenderung kecil, sedangkan untuk kata yang semakin jarang muncul akan memiliki nilai yang cenderung besar.

a. TF (Term Frequency)

Term Frequency (TF) mengembalikan frekuensi *term* (t) di setiap dokumen (d) dari kosakata yang telah diproses sebelumnya frekuensinya (M. T. H. K. Tusar and M. T. Islam 2021) [6]. Atau secara sederhana TF ini akan menghitung frekuensi kemunculan kata pada sebuah dokumen. Umumnya nilai TF ini dibagi jumlah seluruh kata pada dokumen.



$$tf_{t,d} = \frac{n_{t,d}}{\text{Total number of terms (t) in the document (d)}} \quad [6]$$

Ket:

tf = Frekuensi kemunculan kata pada sebuah dokumen.

n = Berapa kali *term* (t) ditemukan dalam dokumen (d).

b. *IDF (Inverse Document Frequency)*

Inverse Document Frequency (IDF) akan menghitung bobot kata-kata penting yang muncul di semua dokumen (M. T. H. K. Tusar and M. T. Islam 2021) [6].

$$idf_t = \log \frac{N}{df_t} \quad [6]$$

Ket:

N = Jumlah dokumen.

df = Jumlah dokumen yang mengandung term (t).

c. *TF-IDF*

Nilai TF dan nilai IDF yang telah diperoleh akan dikalikan untuk mendapatkan nilai TF-IDF.

$$(tf - idf)_{t,d} = tf_{t,d} \times idf_t \quad [6]$$

Berikut adalah beberapa penelitian yang telah dilakukan dengan memanfaatkan metode TF-IDF:

Imamah and F. H. Rachman (2020) [28] melakukan penelitian yang mengumpulkan ulasan data dari Twitter untuk mengetahui opini publik di seluruh dunia tentang COVID-19. Penelitian ini mengusulkan untuk melakukan analisis sentimen agar dapat mengetahui tentang kesehatan mental melalui opini publik di Twitter. Dengan demikian bisa dilihat kesehatan mental masyarakat di dunia, apakah masih merasa aman atau dalam tahap kekhawatiran berlebihan. Dataset yang digunakan dalam penelitian ini adalah data *tweet* COVID-19 yang dikumpulkan pada tanggal 30 April 2020, yang terdiri dari 355384 *tweets review*. Adapun metode yang digunakan adalah *Term Weighting TF-IDF* dan *Logistic Regression* untuk mengklasifikasikan sentimen positif, netral, dan negatif. Hasil

model yang menggunakan algoritma *Logistic Regression* menunjukkan akurasi mencapai 94,71%.



S. Singh et al. (2022) [29] melakukan penelitian yang membahas tentang metode berbasis *Machine Learning* untuk menganalisis sentimen pada *tweet* dari *dataset* Twitter. Penelitian analisis sentimen ini dilakukan pada kumpulan data Twitter yang tersedia secara umum untuk *US Airlines*. Metode ekstraksi fitur dilakukan dengan menggunakan teknik *Term Frequency* dan *Inverse Document Frequency* (TF-IDF), kemudian hasil ekstraksi fitur diberikan sebagai input ke beberapa algoritma klasifikasi yang digunakan secara individual pada penelitian ini, yaitu *random forest*, SVM, *gradient boosting*, dan *extreme gradient boosting* (XGBoost). Hasilnya menunjukkan bahwa algoritma SVM memberikan hasil terbaik dengan tingkat akurasi 83.74%, presisi 84%, dan skor F1 sebesar 87.85%. Adapun tiga metode lain yaitu RF, GB dan XGBoost, masing-masing memperoleh tingkat akurasi sebesar 77.90%, 80.46%, dan 81.55%.

2.1.10 Metaheuristic dan PSO (*Particle Swarm Optimization*)

Algoritma metaheuristik adalah suatu metode optimasi yang dapat mendekati solusi optimal dari berbagai masalah optimasi. Algoritma ini dikenal karena sifatnya yang tidak memerlukan turunan dan memiliki kesederhanaan, fleksibilitas, serta kemampuan untuk menghindari jebakan lokal (local optima). Algoritma metaheuristik memiliki sifat stokastik, dimulai dengan menghasilkan solusi secara acak tanpa perlu menghitung turunan di ruang pencarian, seperti yang dilakukan dalam teknik pencarian gradien. Kelebihan lainnya adalah fleksibilitas dan kemudahan implementasinya, yang memungkinkan algoritma ini dapat dengan mudah dimodifikasi sesuai dengan permasalahan yang sedang dihadapi (P. Agrawal et al. 2021) [41].

Salah satu algoritma yang termasuk dalam metaheuristik ini adalah *Particle Swarm Optimization* (PSO) yang dikembangkan oleh Kennedy dan Eberhart. Algoritma ini memiliki keunggulan dimana implementasinya yang sederhana dan parameter pengontrol yang lebih sedikit. Ide dan konsep dasar algoritma PSO muncul berdasarkan pengamatan perilaku sosial pada kelompok burung dan ikan. Dalam alam, sekelompok burung terbang bersama mengikuti pemimpin yang aling dekat dengan sumber makanan. Konsep ini diadaptasi ke dalam algoritma, khususnya dalam algoritma PSO, untuk menyelesaikan masalah. Dalam konteks PSO, sekelompok partikel dianggap sebagai analogi dari



segerombolan burung, dan setiap partikel mewakili solusi potensial. Partikel-partikel ini bergerak dalam ruang multidimensi untuk menemukan solusi terbaik yang mengoptimalkan permasalahan yang sedang dihadapi (T. M. Shami et al. 2022) [42].

Dalam standard PSO (SPSO), sekelompok partikel bergerak dalam ruang pencarian berdimensi D dengan tujuan mencari solusi yang optimal. Setiap partikel i dalam algoritma SPSO memiliki vektor kecepatan arus $V_i = [v_{i1}, v_{i2}, \dots, v_{iD}]$ dan vektor posisi arus $X_i = [x_{i1}, x_{i2}, \dots, x_{iD}]$, di mana D adalah banyaknya dimensi. Proses SPSO dimulai dengan menginisialisasi V_i dan X_i secara acak. Pada setiap iterasi, posisi terbaik yang ditemukan oleh partikel i , yang disebut $Pbest_i = [Pbest_{i1}, Pbest_{i2}, \dots, Pbest_{iD}]$, dan posisi terbaik yang ditemukan oleh seluruh gerombolan, yang disebut $Gbest = [Gbest_1, Gbest_2, \dots, Gbest_D]$, memandu partikel i untuk memperbarui kecepatan dan posisinya dengan menggunakan rumus berikut (T. M. Shami et al. 2022) [42]:

- a. Memperbarui kecepatan partikel i :

$$v_{id}(t+1) = v_{id}(t) + c_1 r_1 (Pbest_{id}(t) - x_{id}(t)) + c_2 r_2 (Gbest_d(t) - x_{id}(t)) \quad [42]$$

- b. Memperbarui posisi partikel i :

$$x_{id}(t+1) = x_{id}(t) + v_{id}(t+1) \quad [42]$$

Dimana t merujuk pada iterasi ke- t dalam algoritma, c_1 dan c_2 adalah koefisien percepatan kognitif dan sosial, serta r_1 dan r_2 adalah dua nilai acak seragam yang dihasilkan dalam interval $[0, 1]$.

2.2 Metode Penyelesaian Masalah

2.2.1 State of the art penelitian

Tabel 1 State of the art penelitian

No.	Judul Karya Ilmiah, Nama, Tahun Terbit Dan Penerbit	Objek dan Permasalahan	Metode Penyelesaian	Kinerja
	Judul: Analisis Sentimen Tweet	Objek: Tweet publik terhadap pembangunan Ibu Kota Negara baru Masalah: Bagaimana menganalisis sentimen dari opini masyarakat Indonesia terkait	Untuk metode heterogeneous ensemble learning menggunakan metode TF-IDF, Naïve Bayes, SVM,	Kinerja topik yang diusulkan ini diharapkan mampu meningkatkan tingkat akurasi dari proses analisis sentimen dibandingkan ketika menggunakan metode individual learning.



	<i>Homogeneous Ensemble Learning Terhadap Pembangunan Ibu Kota Negara Indonesia</i>	pembangunan Ibu Kota Negara baru.	Logistic Regression, dan weighted soft voting • Untuk metode homogeneous ensemble learning menggunakan metode TF-IDF, Bagging dan weighted soft voting	
2	Judul: A Comparative Sentiment Analysis Of Sentence Embedding Using Machine Learning Techniques Penulis: A. Poornima and K. S. Priya [14] Tahun: 2020 Penerbit: IEEE	Objek: Tweet opini teks tentang produk atau acara Permasalahan: Bagaimana membandingkan kinerja algoritma pembelajaran mesin yang berbeda dalam melakukan analisis sentimen data Twitter?	Logistic regression, SVM, dan Naïve Bayes	• Logistic Regression = 86.23% • SVM = 85.69% • Naïve Bayes = 83.54%
3	Judul: A Comparative Study of Sentiment Analysis Using NLP and Different Machine Learning Techniques on Us Airline Twitter Data Penulis: M. T. H. K. Tusar and M. T. Islam [6] Tahun: 2021 Penerbit: IEEE	Objek: US Airline Twitter Data Permasalahan: Dibutuhkannya pendekatan terbaik dalam menganalisis sentimen pada dataset Twitter US Airline	BoW dan TF-IDF untuk NLP, serta SVM, Naïve Bayes, Random Forest, dan Logistic Regression untuk machine learning algorithm	SVM dan Logistic Regression memperoleh akurasi tertinggi 77%, ketika BoW dan TF-IDF digabungkan dengan setiap algoritma machine learning.
	Judul: Sentiment Analysis in Twitter Using Lexicon Polarity ion Zusrini shuri [2]	Objek: Tweet umum dari Twitter Permasalahan: Bagaimana melakukan analisis sentimen menggunakan metode unsupervised learning	Lexicon dan Polarity Multiplication	Tingkat akurasi mencapai 68%



	Tahun: 2019 Penerbit: IEEE			
5	Judul: Twitter Sentiments Analysis Using Machine Learning Methods Penulis: L. Mandloi and R. Patel [3] Tahun: 2020 Penerbit: IEEE	Objek: Tweet umum dari Twitter Permasalahan: Bagaimana melakukan analisis sentimen dari setiap tweet masyarakat menggunakan algoritma machine learning dan membandingkan performanya	Naïve bayes, SVM, dan Maximum Entropy	- Tingkat akurasi dan presisi Naïve Bayes mencapai 86% dan 88.7% - Tingkat akurasi dan presisi Maximum Entropy mencapai 82.6% dan 84.05% - Tingkat akurasi dan presisi SVM mencapai 74.6% dan 75.9%
6	Judul: Sentiment Analysis of Twitter Data Penulis: S. A. El Rahman et al. [15] Tahun: 2019 Penerbit: IEEE	Objek: Tweet ulasan tentang McDonalds dan KFC Permasalahan: Data Twitter yang tidak terstruktur dan sulit untuk dianalisis, terkhususnya dalam menganalisis ulasan masyarakat terkait McDonalds dan KFC	Model berbasis lexicon untuk unsupervised learning, serta Naïve bayes, SVM, maximum entropy, decision tree, random forest, dan bagging untuk supervised learning	- Untuk data McDonalds memperoleh skor akurasi cross validate 74%. - Untuk data KFC memperoleh skor akurasi cross validate 78%.
7	Judul: Analysing The Twitter Sentiments in Covid-19 Using Machine Learning Algorithms Penulis: R. Katarya et al. [16] Tahun: 2022 Penerbit: IEEE	Objek: Opini masyarakat tentang penyebaran COVID-19 Permasalahan: Bagaimana memproses opini masyarakat terkait penyebaran COVID-19	Logistic Regression, Multinomial Naive Bayes, KNN, Random Forest, dan SVM	- KNN memperoleh skor presisi tertinggi sebesar 81% - Logistic regression memperoleh skor recall, F1, akurasi, dan balanced accuracy tertinggi sekitar 73%.
8	Judul: Sentiment Analysis of COVID-19 Vaccines from Indonesian Tweets Penulis: Using machine learning Tahun: 2022 Penerbit: IEEE	Objek: Postingan online media sosial tentang vaksin COVID-19 di Indonesia Permasalahan: Peningkatan berbagai macam respon masyarakat Indonesia di media sosial	SVM, Logistic Regression, KNN, Multinomial Naïve Bayes, Random Forest, Decision Tree, dan Multilayer Perceptron (MLP)	Ketika menggunakan BoW algoritma Naïve Bayes memiliki skor F1 tertinggi sebesar 70%, dan akurasi tertinggi dimiliki oleh algoritma SVM, Logistic Regression, Naïve Bayes, dan



	Penulis: R. Latifah et al. [17] Tahun: 2021 Penerbit: IEEE	mengenai vaksin COVID-19		MLP di kisaran 74-76% Ketika menggunakan TF-IDF algoritma SVM dan logistic regression memiliki akurasi dan skor f1 tertinggi, masing-masing 84% dan 83%, serta skor F1 masing-masing 76%.
9	Judul: Comparative Study of Twitter Sentiment on COVID - 19 Tweets Penulis: A. J. Nair et al. [19] Tahun: 2021 Penerbit: IEEE	Objek: Tweet perkembangan COVID-19. Permasalahan: Perkembangan jumlah tweet mengenai COVID-19 di Twitter yang mengalami peningkatan pada tingkat yang belum pernah terjadi sebelumnya.	Logistic Regression, VADER, dan BERT	Akurasi: Logistic Regression = 83% BERT = 92% VADER = 88%
10	Judul: Comparative Study of Covid-19 Tweets Sentiment Classification Methods Penulis: U. N. Wisesty et al. [20] Tahun: 2021 Penerbit: IEEE	Objek: Komentar masyarakat di Twitter tentang hal-hal yang terjadi selama pandemi COVID-19 Permasalahan: Populernya topik diskusi terkait covid pada 2 tahun terakhir di berbagai media sosial dan mengundang banyak tanggapan dan opini masyarakat	- Vector space model (BoW dan TF-IDF) digabungkan dengan SVM - Word Embedding (Word2Vec dan Transformers) digabungkan dengan LSTM. - BERT	BERT memperoleh weighted F1-score tertinggi sebesar 0.85 dan macro F1-score sebesar 0.83
11	Judul: Hybrid Method for Sentiment Analysis Using Homogeneous Ensemble Classifier Penulis: Murni et	Objek: Ulasan TripAdvisor (Web ulasan wisata) Permasalahan: Ditemukan adanya ketidakkonsistenan antara komentar wisatawan terhadap objek wisata dengan ulasan atau penilaian mereka yang digambarkan dengan peringkat bintang	- Metode lexical menggunakan SenticNet - Metode machine learning menggunakan algoritma bagged decision tree, logistic model tree, dan bagged multi-layer perceptron	Tingkat akurasi mencapai 98.1267%



12	Judul: Tweet Sentiment Analysis for 2019 Indonesia Presidential Election Results Using Various Classification Algorithms Penulis: F. A. Wenando et al. [4] Tahun: 2020 Penerbit: IEEE	Objek: Hasil pemilu presiden 2019 di Indonesia Permasalahan: Setelah Pilkada berakhir banyak masyarakat yang mengungkapkan opini dan pernyataannya di media sosial mengenai pemilu 2019	• Naïve Bayes, SMO, Logistic Regression, KNN, dan Decision Tree	SMO memperoleh hasil terbaik dari segi akurasi, presisi, dan recall, yaitu masing-masing 82.7%.
13	Judul: Sentiment Analysis Algorithm Based on BERT and Convolutional Neural Network Penulis: R. Man and K. Lin [30] Tahun: 2021 Penerbit: IEEE	Objek: Ulasan hotel (ChnSentiCorp) Permasalahan: Dalam analisis sentimen, model Word2vec tradisional dianggap tidak dapat sepenuhnya mengungkapkan informasi yang terkandung dalam kata-kata.	BERT dan CNN	Akurasi = 91% Recall = 90% F1-score = 90%
14	Judul: Detecting and Understanding Sentiment Trends and Emotion Patterns of Twitter Users - A Study on the Demise of a Bollywood Celebrity Penulis: A. A. Marouf et al. [31] Tahun: 2022 Penerbit: MDPI	Objek: Postingan media sosial dari penggemar aktor Bollywood Sushant Singh Rajput mengenai kasus bunuh dirinya Permasalahan: Bagaimana mendeteksi dan memahami tren sentimen dan pola emosional pengguna media sosial dalam kasus bunuh diri selebriti	• Untuk menganalisis sentimen tweet digunakan NLTK dan Textblob Package. • Untuk mengekstrak fitur psikolinguistik, seperti skor emosi positif dan negatif, dari teks digunakan LIWC (Linguistic Inquiry and Word Count)	• Sentimen positif tertinggi (51,2%) dan emosi positif (63,1%) dilaporkan untuk tagar #Plant4SSR. • Sentimen negatif tertinggi (38,9%) dilaporkan untuk #SSR dan emosi negatif tertinggi (70,2%) dilaporkan untuk #JusticeForSSR
	 -Based r	Objek: Ulasan objek wisata	ELECTRA-Based Pipeline	Model yang diusulkan sangat efisien dalam analisis sentimen ulasan objek wisata dalam hal

	Sentiment Analysis of Tourist Attraction Reviews Penulis: H. Fang et al. [32] Tahun: 2022 Penerbit: MDPI	Permasalahan: Ulasan objek wisata yang bersifat informal dan gaduh, serta sebagian besar penelitian yang menganalisisnya hanya berfokus pada shallow machine learning models atau model deep learning yang tidak terlatih		kinerja model, Ini mewakili, rata-rata, peningkatan 6.33% dalam m_{avgF} dan rata-rata peningkatan 1.5% pada w_{avgF}.
16	Judul: A Sentiment Analysis Approach to Predict an Individual's Awareness of the Precautionary Procedures to Prevent COVID-19 Outbreaks in Saudi Arabia Penulis: S. S. Aljameel et al [33] Tahun: 2020 Penerbit: MDPI	Objek: Tweet terkait COVID-19 di Arab Saudi Permasalahan: Dalam situasi pandemi COVID-19 perlunya keputusan cepat dibuat mengenai praktik yang harus diikuti negara untuk mengendalikan penyakit dan prosedur pencegahan yang sesuai untuk setiap wilayah berdasarkan tingkat kesadarannya.	SVM, KNN, dan Naïve Bayes	Bigram TF-IDF dengan classifier SVM menghasilkan akurasi tertinggi sebesar 85%
17	Judul: Comparison of Different Modeling Techniques for Flemish Twitter Sentiment Analysis Penulis: M. Reusens et al. [34] Tahun: 2022 Penerbit: MDPI	Objek: Tweet bahasa Flemish Permasalahan: Penelitian analisis sentimen meningkat popularitasnya selama beberapa dekade terakhir, namun sebagian besar penelitian berfokus pada bahasa Inggris, yang membuat bahasa lain kurang terwakili	Lexicon-Based Method, Traditional Machine-Learning Models, Neural Network, dan Attention-Based Model	• Hasil Lexicon-Based Method terbaik adalah model TextBlob dengan akurasi 43% • Hasil metode Traditional Machine-Learning terbaik adalah model Logistic Regression dengan akurasi 61% • Hasil metode Neural Network terbaik adalah model Bidirectional LSTM dengan akurasi 61% • Hasil metode Attention-Based terbaik adalah model Fine-tuned RobBERT dengan akurasi 66%



18	Judul: A Semi-Supervised Approach to Sentiment Analysis of Tweets during the 2022 Philippine Presidential Election Penulis: J. J. E. Macrohon et al. [35] Tahun: 2022 Penerbit: MDPI	Objek: Tweet selama pemilihan presiden Filipina 2022 Permasalahan: Bagaimana mendapatkan sentimen umum orang Filipina selama Pemilihan Presiden Filipina 2022.	• Word Embedding menggunakan TF-IDF • Hyperparameter Tuning menggunakan GridSearchCV • Semi-Supervised Learning menggunakan Multinomial Naïve Bayes	Tingkat akurasi mencapai 84.83%
----	--	--	---	--

Berdasarkan uraian dari Tabel 1 *State of The Art* di atas, diperoleh hasil mengenai tantangan dan peluang untuk penelitian ini, yang berkaitan dengan analisis sentimen dan *ensemble learning*, yaitu:

- Penelitian [2] masih menggunakan metode *unsupervised learning*, yaitu lexicon untuk melakukan analisis sentimen, yang dimana metode ini memiliki tingkat akurasi yang rendah.
- Penelitian [2][3][4][6][14][15][16][17][19][20][31][32][33][34][35] masih melakukan analisis sentimen menggunakan metode *individual learning*, yaitu hanya menggunakan sebuah algoritma *machine learning* untuk melakukan analisis sentimen dan membandingkan setiap algoritma. Penggunaan metode *individual learning* ini tentu akan memiliki kekurangan tergantung pada setiap algoritma yang digunakan, dan tentunya hal tersebut dapat diatasi dengan menggabungkan beberapa algoritma (*ensemble*) agar setiap algoritma yang digabungkan dapat menutupi kekurangan algoritma yang lain.
- Penelitian [31][32] tidak menghitung tingkat akurasi dari analisis sentimennya, sehingga pada penelitian ini akan dihitung tingkat akurasi dari analisis sentimen menggunakan metode *heterogeneous ensemble learning* dan *homogeneous ensemble learning*, serta membandingkan keduanya.



Penelitian [26][30] mengusulkan metode gabungan beberapa algoritma *machine learning* yang berbeda untuk melakukan analisis sentimen dengan tingkat akurasi yang cukup baik tetapi masih menggunakan metode *hard voting* untuk melakukan klasifikasi *binary*.

2.2.2 Metode yang diusulkan

Berdasarkan *State of The Art* di atas, maka pada penelitian ini akan diusulkan model *supervised learning*, dimana akan diterapkan dua metode *Ensemble Learning*, yaitu *Heterogeneous* dan *Homogeneous Ensemble Learning* dalam melakukan analisis sentimen multikelas dan akan membandingkan kinerja keduanya. Tiga algoritma *machine learning* akan diimplementasikan pada model *ensemble learning* yang akan dibangun, yaitu *Naïve Bayes*, *SVM*, dan *Logistic Regression*. Ketiga algoritma ini dipilih karena telah banyak digunakan dalam proses analisis sentimen *individual learning*, selain itu ketiganya merupakan algoritma *machine learning* yang sederhana tapi cepat, dan memiliki tingkat akurasi yang cukup baik dalam melakukan proses klasifikasi. Meskipun demikian, tingkat akurasi dari penerapan *individual learning* pada setiap algoritma tersebut masih dapat ditingkatkan dengan menggunakan metode *ensemble learning*, meskipun akan mengorbankan sedikit kecepatan komputasinya.

Untuk metode *heterogeneous ensemble learning* sendiri, algoritma *Naïve Bayes*, *SVM*, dan *Logistic Regression* akan digabungkan. Tiga algoritma yang digabungkan ini akan dilatih dan diuji masing-masing dengan data latih dan data uji yang sama, kemudian hasil pengujian setiap algoritma akan di-*voting* untuk mendapatkan final klasifikasinya. Penerapan tiga algoritma klasifikasi *machine learning* *Naïve Bayes*, *SVM*, dan *Logistic Regression* ke dalam metode *Heterogeneous Ensemble Learning* yang diusulkan ini diharapkan dapat memberikan peningkatan skor akurasi dari proses analisis sentimen, dibandingkan ketika menggunakan metode *individual learning*.

Kemudian untuk metode *homogeneous ensemble learning* akan diimplementasikan menggunakan metode *Bagging (Bootstrap Aggregating)* dalam proses *ensemble*-nya, agar dapat mengurangi *varians sampling* seperti *overfitting* serta dapat terbentuk sebuah model baru yang lebih kokoh dan akurat ketika bekerja dengan *dataset* yang besar. Pada metode ini juga akan menggunakan 3 algoritma *machine learning* seperti yang diterapkan pada metode *heterogeneous*, yaitu *Naïve*

VM, dan *Logistic Regression*. Akan tetapi untuk metode *homogeneous* usulkan ini, ketiga algoritma tersebut akan diimplementasikan kedalam tiga dimana pada skenario pertama akan menggunakan algoritma *Naïve Bayes*



sebagai *base classifier*-nya, skenario kedua akan menggunakan SVM sebagai *base classifier*-nya, dan skenario ketiga akan menggunakan *Logistic Regression* sebagai *base classifier*-nya, sehingga nantinya dapat dilihat *base classifier* mana yang akan memberikan skor akurasi terbaik ketika dimasukkan ke dalam metode *homogeneous ensemble learning*. Data latih dan data uji yang sama juga akan digunakan pada setiap algoritma di setiap skenario, kemudian hasil pengujian setiap algoritma di setiap skenario akan di-*voting* untuk mendapatkan final klasifikasi dari setiap skenarionya. Dengan menerapkan metode *homogeneous ensemble learning* pada ketiga algoritma tersebut, diharapkan dapat memberikan peningkatan akurasi analisis sentimen dibandingkan ketika menggunakan algoritma-algoritma tersebut secara individual.

Terakhir, karena yang akan dibangun adalah metode *ensemble learning*, tentunya akan menerapkan proses *voting* di akhir untuk menentukan final klasifikasi. Kebanyakan metode *ensemble learning* menerapkan *hard voting* dan *simple soft voting* untuk proses *voting* final klasifikasinya, akan tetapi pada metode *ensemble learning* yang diusulkan dalam penelitian ini (*Heterogeneous* dan *Homogeneous*) akan menggunakan metode *weighted soft voting*. Metode *weighted soft voting* ini dipilih karena menerapkan keunggulan dari metode *simple soft voting* yang dikombinasikan dengan penambahan bobot (*weight*) dari setiap *individual classifier* yang digabungkan. Metode *simple soft voting* memiliki keunggulan dalam konsep penggabungan nilai probabilitas yang dapat mengatasi ketidakseimbangan kelas pada *dataset*, dimana terkadang terdapat ketidakseimbangan antara jumlah sampel pada setiap kelas. *Simple soft voting* dapat mengatasi hal tersebut dengan menggabungkan probabilitas dari setiap model, dimana setiap model akan memberikan kontribusi berdasarkan probabilitas prediksi kelasnya, sehingga dapat memberikan bobot yang lebih seimbang pada setiap kelas, termasuk kelas minoritas, dan hal ini juga bisa sangat berguna dalam klasifikasi multikelas.

Penambahan bobot (*weight*) dari setiap *individual classifier* ke dalam metode *simple soft voting* akan dilakukan untuk membangun metode *weighted soft voting*



in digunakan dalam penelitian ini. Dalam proses *voting* yang akan ing pada metode *simple soft voting*, faktor bobot dari masing-masing ng dihasilkan oleh *individual classifier* akan menjadi pertimbangan. Nilai

bobot setiap *individual classifier* ditentukan menggunakan algoritma PSO, dimana kombinasi nilai bobot yang menghasilkan performa terbaik pada model *ensemble* secara keseluruhan akan dipilih untuk diterapkan. Hal ini akan memberikan nilai bobot (*weight*) yang optimal pada setiap *individual classifier*, yang kemudian akan dimasukkan ke dalam *soft voting*, sehingga dapat menghasilkan hasil yang lebih baik secara keseluruhan untuk model *ensemble* yang dibangun. Selain itu, karena metode *weighted soft voting* ini membutuhkan nilai probabilitas untuk proses *voting*-nya, maka hasil keluaran setiap model dari setiap algoritma bukan berupa kelas hasil klasifikasi/prediksi, melainkan berupa nilai probabilitas dari setiap kelas. Kemudian final klasifikasi akan diambil berdasarkan kelas dengan nilai probabilitas tertinggi yang dirata-ratakan atas semua *individual classifier* dengan tetap mempertimbangkan bobotnya.

Seperti yang dijelaskan sebelumnya, metode *Particle Swarm Optimization* (PSO) akan digunakan dalam menentukan kombinasi nilai bobot yang optimal untuk setiap model. Dalam proses pencarian kombinasi bobot yang optimal untuk setiap *trained model*, partikel-partikel dalam PSO akan saling berkomunikasi, bergerak bersama, dan saling berkoordinasi untuk mencapai performa yang optimal. Dengan kata lain, algoritma PSO akan berusaha untuk menemukan kombinasi bobot yang dapat menghasilkan performa terbaik secara keseluruhan untuk model *ensemble* yang dibangun.

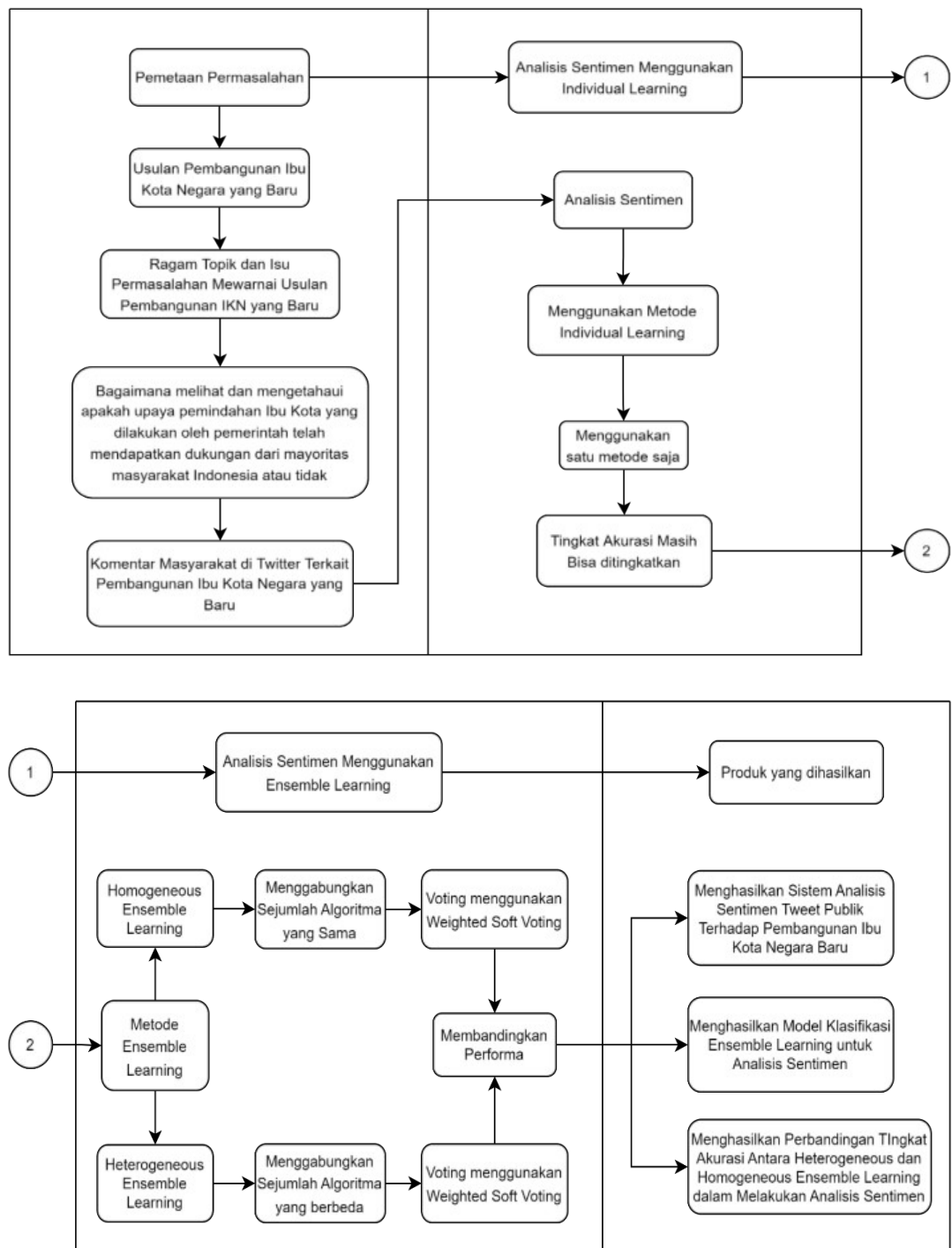
2.3 Target Hipotesis Hasil Penelitian

Berdasarkan Tabel 1 *State of The Art* dan penelitian yang akan dikerjakan, maka penelitian ini menargetkan mampu menghasilkan sebuah model analisis sentimen menggunakan metode *ensemble learning* dengan tingkat akurasi yang lebih baik dari metode *individual learning* dalam melakukan analisis sentimen *tweet* terkait pembangunan Ibu Kota Negara baru Indonesia. Selain itu, penelitian ini juga diharapkan mampu memberikan hasil perbandingan yang jelas antara metode *heterogeneous ensemble learning* dan metode *homogeneous ensemble learning*



melakukan analisis sentimen, sehingga nantinya dapat diketahui metode yang memiliki performa terbaik dalam melakukan analisis sentimen *tweet* terhadap pembangunan Ibu Kota Negara baru Indonesia.

2.4 Kerangka Pikir Penelitian



Gambar 4 Kerangka pikir penelitian

