

**ANALISIS PERBANDINGAN KEMIRIPAN TEKS BAHASA
DAERAH DI INDONESIA MENGGUNAKAN ALGORITMA
RANDOM FOREST DAN *SUPPORT VECTOR MACHINE***

SKRIPSI

Program Studi Sarjana Informatika
Jurusan Informatika

Oleh:
FATHIMAH NUR UMMAMI
NIM D1041201063



FAKULTAS TEKNIK
UNIVERSITAS TANJUNGPURA
PONTIANAK
2024

HALAMAN PERNYATAAN

Yang bertanda tangan di bawah ini:

Nama : Fathimah Nur Ummami

NIM : D1041201063

menyatakan bahwa dalam skripsi yang berjudul “Analisis Perbandingan Kemiripan Teks Bahasa Daerah di Indonesia Menggunakan Algoritma *Random Forest* dan *Support Vector Machine*” tidak terdapat karya yang pernah diajukan untuk memperoleh gelar sarjana di suatu perguruan tinggi manapun. Sepanjang pengetahuan Saya, tidak terdapat karya atau pendapat yang pernah ditulis atau diterbitkan oleh orang lain, kecuali yang secara tertulis diacu dalam naskah ini dan disebutkan dalam Daftar Pustaka.

Demikian pernyataan ini dibuat dengan sebenar-benarnya. Saya sanggup menerima konsekuensi akademis dan hukum di kemudian hari apabila pernyataan yang dibuat ini tidak benar.

Pontianak, 5 Desember 2024

Fathimah Nur Ummami
NIM D1041201063



KEMENTERIAN PENDIDIKAN, KEBUDAYAAN,
RISET, DAN TEKNOLOGI
UNIVERSITAS TANJUNGPURA
FAKULTAS TEKNIK

Jalan Prof. Dr. H. Hadari Nawawi Pontianak 78124 Telp. (0561) 740186
Kotak Pos 1049 Email: ft@untan.ac.id Website: http://Teknik.untan.ac.id

HALAMAN PENGESAHAN

ANALISIS PERBANDINGAN KEMIRIPAN TEKS BAHASA DAERAH DI
INDONESIA MENGGUNAKAN ALGORITMA
RANDOM FOREST DAN SUPPORT VECTOR MACHINE

Program Studi Sarjana Informatika
Jurusan Informatika

Oleh:

Fathimah Nur Ummami
NIM D1041201063

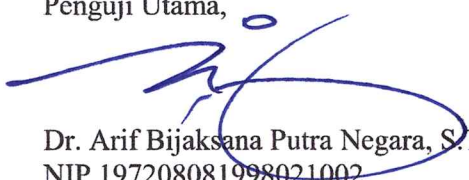
Telah dipertahankan di depan Penguji Skripsi pada tanggal 5 Desember 2024
dan diterima sebagai salah satu persyaratan untuk memperoleh gelar sarjana.

Susunan Penguji Skripsi:

Ketua,


Prof. Dr. Henry Sujaini, S.T., M.T.
NIP 196806291997021001

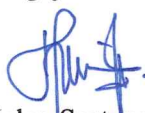
Penguji Utama,


Dr. Arif Bijaksana Putra Negara, S.T., M.T.
NIP 197208081998021002

Sekretaris,


Rina Septiriana, S.T., M.Cs.
NIP 198709232020122001

Penguji Pendamping,


Helen Sastypratiwi, S.T., M.Eng.
NIP 198601172012122004

Pontianak, 5 Desember 2024

Dekan,


Dr.-Ing. Ir. Slamet Widodo, M.T., IPM
NIP 196712231992031002

HALAMAN PERSEMBAHAN

Terima kasih untuk diri sendiri yang telah memulai apa yang sudah dipilih dan menyelesaikannya dengan baik. Tentu tidak mudah untuk sampai dititik ini, mungkin langkah ini tidak sepenuhnya mendapatkan beberapa peran di kehidupan layaknya seperti anak yang lain, namun langkah ini tetap terus berjalan dan kuat karena diiringi dengan doa-doa tulus dari mama. Skripsi ini saya persembahkan kepada orang-orang yang sangat saya sayangi dan cintai.

Teruntuk mama dan ayah, dengan segala hormat cinta, kasih sayang, dan terima kasih saya ucapkan yang sebesar-besarnya. Saya persembahkan skripsi ini sebagai tanda tanggung jawab dalam menjalani apa yang sudah menjadi pilihan menuju jalan masa depan. Terima kasih selalu memberikan doa yang tak pernah putus, pengorbanan, dan terima kasih atas pelajaran hidup yang mungkin bisa jadi bekal untuk kedepannya.

Terima kasih penulis ucapkan kepada sahabat-sahabat yang senantiasa selalu menemani dan memberikan support tiada habisnya yaitu Sherina Tri Manggiri, Fiyen Febryan, Evelyn Oktaviana Hutapea, Salsabil Humaira, Tari Wahyuni, Cik Dara Kamila, Jesica Fitriani Silalahi, Putri Sintang Anugeraeni, dan Tri Rizky Saputri.

Teruntuk pembimbing akademik selama masa perkuliahan yaitu ibu Helen Sastypratiwi, S.T., M.Eng., dengan segala hormat dan terima kasih saya ucapkan kepada ibu yang sudah memberikan bimbingan, arahan, dan ilmu yang berharga kepada saya selama masa perkuliahan.

Teruntuk pembimbing tugas akhir, dengan segala hormat dan terima kasih sebesar-besarnya saya ucapkan kepada Bapak Prof. Dr. Herry Sujaini, S.T., M.T, dan ibu Rina Septiriana, S.T., M.Cs. Terima kasih kepada bapak dan ibu yang telah memberikan segala bantuan, bimbingan, nasihat, ilmu, dan kesabaran yang telah bapak dan ibu berikan selama menjadi pembimbing saya yang akan selalu saya ingat. Saya mohon maaf apabila ada kesalahan yang saya perbuat selama masa bimbingan.

Terima kasih juga kepada teman seperjuangan di masa menempuh perkuliahan yaitu Fika Tsalsabila Tinanda dan Rendy Kurniawan, serta terima kasih kepada teman-teman seperjuangan INFORMATIKA 20 dan teman-teman SETARA TEKNIK 20.

KATA PENGANTAR

Puji dan Syukur penulis panjatkan atas kehadiran Allah SWT, yang telah memberikan Rahmat serta karunia-Nya, Sholawat serta salam yang senantiasa dilimpahkan kepada junjungan Nabi Muhammad SAW, sehingga penulis dapat menyelesaikan skripsi ini dengan judul “Analisis Perbandingan Kemiripan Teks Bahasa Daerah di Indonesia Menggunakan Algoritma *Random Forest* dan *Support Vector Machine*”. Sebagai syarat untuk menyelesaikan Program Sarjana (S1) pada Fakultas Teknik, Jurusan Informatika, Universitas Tanjungpura.

Penulis menyadari bahwa dalam penulisan skripsi ini tidak terlepas dari bantuan berbagai pihak, baik dari segi moril dan materil. Oleh karena itu, penulis menyampaikan rasa hormat dan terima kasih yang sebesar-besarnya kepada pihak-pihak yang telah turut andil memberikan bantuan dan dukungan kepada penulis. Secara khusus penulis menyampaikan ucapan terima kasih kepada Bapak Prof. Dr. Herry Sujaini, S.T., M.T. selaku dosen Pembimbing I dan Ibu Rina Septiriana, S.T., M.Cs. selaku dosen Pembimbing II, serta kepada Bapak Dr. Arif Bijaksana Putra Negara, S.T., M.T. selaku dosen Penguji I dan Ibu Helen Sastypratiwi, S.T., M.Eng. selaku dosen Penguji II.

Penulis menyadari bahwa skripsi ini belum sempurna dan masih terdapat kekurangan. Oleh karena itu, dengan segenap kerendahan hati penulis mengharapkan saran dan kritik yang membangun dari semua pihak untuk perbaikan dan pengembangan penulisan berikutnya. Semoga hasil penelitian ini dapat memberikan kontribusi yang bermanfaat bagi pembelajaran dan pemahaman mendalam, terutama mengenai Analisis Perbandingan Kemiripan Teks Bahasa Daerah di Indonesia Menggunakan Algoritma *Random Forest* dan *Support Vector Machine*.

Pontianak, 5 Desember 2024
Penulis,

Fathimah Nur Ummami

ABSTRAK

Indonesia memiliki keberagaman bahasa daerah, termasuk di Kalimantan Barat yang terdiri dari beberapa bahasa daerah seperti bahasa Dayak Iban, Dayak Pesaguan, Melayu Pontianak, dan Melayu Sambas. Kemiripan bahasa dalam satu rumpun sering kali menyebabkan kesalahan klasifikasi dalam sistem pengenalan bahasa otomatis. Oleh karena itu, penelitian ini dilakukan bertujuan untuk membandingkan performa algoritma *Random Forest* dan *Support Vector Machine* (SVM) dalam mengklasifikasikan teks bahasa daerah serta mengidentifikasi tingkat kemiripan antar bahasa yang digunakan dalam penelitian. Metodologi yang diterapkan pada penelitian meliputi pengumpulan data dari korpus Nusantara, bahasa yang digunakan dalam penelitian ini yaitu bahasa Dayak Iban (997 data), bahasa Dayak Pesaguan (1603 data), bahasa Indonesia (9099 data), bahasa Melayu Pontianak (1747 data), dan bahasa Melayu Sambas (9099 data). Selanjutnya dilakukan pra-pemrosesan teks (*cleaning*, *case folding*, tokenisasi, dan vektorisasi TF-IDF), serta penanganan ketidakseimbangan data menggunakan teknik SMOTE. Model kemudian dikembangkan menggunakan dua pendekatan, yaitu data Non-SMOTE dan data SMOTE, dengan algoritma *Random Forest* dan SVM. Evaluasi model dilakukan dengan menggunakan *heatmap confusion matrix* untuk menganalisis kesalahan klasifikasi serta menghitung *accuracy* model. Hasil penelitian menunjukkan bahwa algoritma SVM (SMOTE) mencapai *accuracy* tertinggi sebesar 95,36%, sementara *Random Forest* (SMOTE) memperoleh *accuracy* sebesar 94,18%. Dari analisis kesalahan klasifikasi, ditemukan bahwa pasangan bahasa Indonesia dan Melayu Sambas memiliki tingkat kemiripan tertinggi sebesar 6,45%. Kesimpulannya, algoritma SVM (SMOTE) lebih baik dibandingkan *Random Forest* dalam mengklasifikasikan bahasa daerah di Kalimantan Barat.

Kata Kunci: Kemiripan Bahasa, *Random Forest*, *Support Vector Machine*, SMOTE, Bahasa Daerah

ABSTRACT

Indonesia has a diverse range of regional languages, including those in West Kalimantan, which consists of several local languages such as Dayak Iban, Dayak Pesaguan, Malay Pontianak, and Malay Sambas. The similarity of languages within the same linguistic family often leads to misclassification in automatic language recognition systems. Therefore, this study aims to compare the performance of the Random Forest and Support Vector Machine (SVM) algorithms in classifying regional language texts and identifying the degree of similarity among the languages used in the research. The methodology applied in this study includes data collection from the Nusantara corpus. The languages used in this study are Dayak Iban (997 data samples), Dayak Pesaguan (1,603 data samples), Indonesian (9,099 data samples), Malay Pontianak (1,747 data samples), and Malay Sambas (9,099 data samples). The next steps involve text preprocessing (cleaning, case folding, tokenization, and TF-IDF vectorization) and handling data imbalance using the SMOTE technique. The model was then developed using two approaches: Non-SMOTE data and SMOTE data, with Random Forest and SVM algorithms. Model evaluation was conducted using a heatmap confusion matrix to analyze classification errors and calculate model accuracy. The results show that the SVM (SMOTE) algorithm achieved the highest accuracy of 95.36%, while Random Forest (SMOTE) obtained an accuracy of 94.18%. From the classification error analysis, it was found that the Indonesian and Malay Sambas language pair had the highest similarity rate of 6.45%. In conclusion, the SVM (SMOTE) algorithm outperforms Random Forest in classifying regional languages in West Kalimantan.

Keywords: *Language Similarity, Random Forest, Support Vector Machine, SMOTE, Regional Languages*

DAFTAR ISI

Halaman Pernyataan	ii
Halaman Pengesahan.....	iii
Halaman Persembahan.....	iv
Kata Pengantar.....	v
Abstrak.....	vi
Abstract.....	vii
Daftar Isi	viii
Daftar Tabel.....	x
Daftar Gambar	xi
Kode Program	xiii
Bab I Pendahuluan.....	1
1.1 Latar Belakang.....	1
1.2 Perumusan Masalah.....	4
1.3 Tujuan Penelitian.....	4
1.4 Pembatasan Masalah	4
1.5 Sistematika Penulisan	4
Bab II Tinjauan Pustaka	1
2.1 Penelitian Terkait.....	6
2.2 Landasan Teori	9
2.2.1 <i>Natural Language Processing (NLP)</i>	9
2.2.2 <i>Term Frequency-Inverse Document Frequency (TF-IDF)</i>	9
2.2.3 Bahasa Daerah.....	9
2.2.4 <i>Synthetic Minority Over-sampling Technique (SMOTE)</i>	10
2.2.5 <i>Random Forest</i>	11
2.2.6 <i>Support Vector Machine (SVM)</i>	12

2.2.7	<i>Confusion Matrix</i>	13
Bab III	Metodologi Penelitian	15
3.1	Metode Penelitian	15
3.2	Pengumpulan Data.....	16
3.3	Persiapan Data	16
3.4	Implementasi Algoritma	18
3.4.1	Random Forest	18
3.4.2	Support Vector Machine (SVM)	20
3.5	Analisis Hasil Implementasi Algoritma	22
3.5.1	Nilai <i>Accuracy</i>	22
3.5.2	<i>Heatmap Confusion Matrix</i>	23
3.6	Evaluasi Hasil	25
Bab IV	Hasil dan Analisis	28
4.1	Hasil Penelitian.....	28
4.2	Hasil Pengumpulan Data	28
4.3	Hasil Persiapan Data.....	31
4.4	Hasil Implementasi Algoritma	37
4.4.1	Random Forest	37
4.4.2	Support Vector Machine (SVM)	40
4.5	Analisis Implementasi Algoritma	44
4.5.1	Nilai <i>Accuracy</i>	44
4.5.2	<i>Heatmap Confusion Matrix</i>	45
4.6	Evaluasi Hasil	70
Bab V	Kesimpulan dan Saran	76
5.1	Kesimpulan.....	76
5.2	Saran.....	76
	Daftar Pustaka	77

DAFTAR TABEL

Tabel 2. 1 Confusion Matrix	13
Tabel 3. 1 Nilai Accuracy Algoritma	22
Tabel 3. 2 Nilai Accuracy Algoritma	25
Tabel 3. 3 Contoh Data Kesalahan Prediksi Pada Pasangan Bahasa	26
Tabel 3. 4 Contoh Kata Dominan Dalam Kesalahan Prediksi	26
Tabel 4. 1 Hasil Pengumpulan Data	29
Tabel 4. 2 Korpus Bahasa Dayak Iban	30
Tabel 4. 3 Korpus Bahasa Dayak Pesaguan	30
Tabel 4. 4 Korpus Bahasa Indonesia	30
Tabel 4. 5 Korpus Bahasa Melayu Pontianak	31
Tabel 4. 6 Kopus Bahasa Melayu Sambas	31
Tabel 4. 7 Nilai Accuracy Algoritma	44
Tabel 4. 8 Nilai Accuracy Algoritma	70
Tabel 4. 9 Data Kesalahan Prediksi Pada "Indonesia" Yang Diprediksi Sebagai "Melayu Sambas" SVM (SMOTE).....	74
Tabel 4. 10 Kata Dominan Dalam Kesalahan Prediksi Pada "Indonesia" Yang Diprediksi Sebagai "Melayu Sambas" SVM (SMOTE)	74

DAFTAR GAMBAR

Gambar 3. 1 Diagram Alir.....	15
Gambar 3. 2 Jumlah Dataset.....	16
Gambar 3. 3 Implementasi Random Forest.....	19
Gambar 3. 4 Implementasi Support Vector Machine (SVM).....	21
Gambar 3. 5 Pasangan Bahasa.....	23
Gambar 3. 6 Contoh Heatmap Confusion Matrix.....	24
Gambar 3. 7 Contoh Sajian Nilai Kemiripan Bahasa	24
Gambar 3. 8 Contoh 3 Pasangan Bahasa Dengan Nilai Kemiripan Tertinggi	25
Gambar 4. 1 Korpus Nusantara	28
Gambar 4. 2 Jumlah Dataset (Korpus)	29
Gambar 4. 3 Data Sebelum Pra-Pemrosesan.....	32
Gambar 4. 4 Hasil Cleaning Teks.....	32
Gambar 4. 5 Hasil Case Folding.....	33
Gambar 4. 6 Hasil Tokenisasi.....	34
Gambar 4. 7 Hasil Vektorisasi Teks (TF-IDF).....	35
Gambar 4. 8 Data Sebelum SMOTE	36
Gambar 4. 9 Data Setelah SMOTE	37
Gambar 4. 10 Hasil Implementasi Random Forest (tanpa SMOTE).....	38
Gambar 4. 11 Hasil Implementasi Random Forest (SMOTE)	38
Gambar 4. 12 Heatmap Confusion Matrix Random Forest (tanpa SMOTE).....	39
Gambar 4. 13 <i>Heatmap Confusion Matrix Random Forest (SMOTE)</i>	40
Gambar 4. 14 Hasil Implementasi Algoritma <i>Support Vector Machine</i> (SVM) (tanpa SMOTE).....	41
Gambar 4. 15 Hasil Implementasi Algoritma <i>Support Vector Machine</i> (SVM) (SMOTE).....	42
Gambar 4. 16 Heatmap Confusion Matrix <i>Support Vector Machine</i> (SVM) (tanpa SMOTE).....	43
Gambar 4. 17 Heatmap Confusion Matrix <i>Support Vector Machine</i> (SVM) (SMOTE).....	44
Gambar 4. 18 Heatmap Confusion Matrix Random Forest (tanpa SMOTE).....	45

Gambar 4. 19 <i>Heatmap Confusion Matrix Random Forest (SMOTE)</i>	50
Gambar 4. 20 <i>Heatmap Confusion Matrix Support Vector Machine (SVM) (tanpa SMOTE)</i>	55
Gambar 4. 21 <i>Heatmap Confusion Matrix Support Vector Machine (SVM) (SMOTE)</i>	60
Gambar 4. 22 <i>Nilai Kemiripan Bahasa RF (Non-SMOTE)</i>	65
Gambar 4. 23 <i>Nilai Kemiripan Bahasa RF (SMOTE)</i>	66
Gambar 4. 24 <i>Nilai Kemiripan Bahasa SVM (Non-SMOTE)</i>	68
Gambar 4. 25 <i>Nilai Kemiripan Bahasa SVM (SMOTE)</i>	69
Gambar 4. 26 <i>Nilai Kemiripan Tertinggi RF (Non-SMOTE)</i>	71
Gambar 4. 27 <i>Nilai Kemiripan Tertinggi RF (SMOTE)</i>	72
Gambar 4. 28 <i>Nilai Kemiripan Tertinggi SVM (Non-SMOTE)</i>	73
Gambar 4. 29 <i>Nilai Kemiripan Tertinggi SVM (SMOTE)</i>	73

KODE PROGRAM

Kode Program 4. 1	Baca Data & Menampilkan Data Sebelum Pra-Pemrosesan	31
Kode Program 4. 2	<i>Cleaning</i> Teks	32
Kode Program 4. 3	<i>Case Folding</i>	33
Kode Program 4. 4	Tokenisasi	33
Kode Program 4. 5	Vektorisasi Teks (TF-IDF)	34
Kode Program 4. 6	Data Sebelum SMOTE	35
Kode Program 4. 7	Data Setelah SMOTE	36
Kode Program 4. 8	Implementasi <i>Random Forest</i>	38
Kode Program 4. 9	<i>Heatmap Confusion Matrix Random Forest</i>	39
Kode Program 4. 10	Implementasi <i>Support Vector Machine</i> (SVM)	41
Kode Program 4. 11	<i>Heatmap Confusion Matrix Support Vector Machine</i> (SVM)	42

BAB I

PENDAHULUAN

1.1 Latar Belakang

Di Indonesia, terdapat 718 bahasa daerah yang digunakan oleh penduduk dari berbagai suku dan etnis yang tersebar di seluruh kepulauan Nusantara. Meskipun lebih dari 90% penduduk Indonesia memahami dan menggunakan Bahasa Indonesia sebagai bahasa komunikasi sehari-hari, bahasa ini bukanlah bahasa ibu bagi sebagian besar masyarakat. Sebaliknya, mayoritas penduduk menggunakan salah satu dari ratusan bahasa daerah tersebut sebagai bahasa ibu mereka. Penggunaan bahasa sehari-hari di Indonesia sering kali bersifat campuran, di mana Bahasa Indonesia yang formal bercampur dengan berbagai dialek Melayu atau bahasa daerah lainnya. Hal tersebut mengakibatkan sulitnya dalam mengenali kemiripan bahasa satu dengan bahasa lainnya (Purnamalia et al., 2023).

Provinsi Kalimantan Barat, wilayah ini memiliki keberagaman budaya dan bahasa, dengan kehadiran berbagai suku seperti Dayak Iban, Dayak Pesaguan, Melayu Pontianak, Melayu Sambas, dan lainnya. Setiap suku memiliki bahasa daerahnya masing-masing yang khas dan mencerminkan identitas budaya mereka. Namun, di tengah keragaman ini terdapat kemiripan antara bahasa-bahasa tersebut, terutama yang berasal dari rumpun bahasa yang sama, seperti Melayu. Namun dari rumpun Dayak juga terdapat kemiripan antara bahasa-bahasa tersebut. Untuk rumpun Melayu adanya kemiripan bahasa antara Melayu Pontianak dengan Melayu Sambas, sedangkan rumpun dayak adanya kemiripan bahasa antara Dayak Iban dengan Dayak Pesaguan. Kemiripan ini sering kali menyebabkan tumpang tindih linguistik, di mana bahasa-bahasa yang digunakan sehari-hari tampak saling berbaur, membuat identifikasi bahasa tertentu menjadi lebih rumit. Kondisi ini semakin diperburuk oleh kebiasaan masyarakat yang secara alami mencampurkan elemen dari berbagai bahasa dalam percakapan sehari-hari, terutama di daerah dengan kontak budaya yang intens antar suku dan etnis.

Dalam konteks ini, klasifikasi bahasa menjadi langkah penting untuk mengelompokkan suatu bahasa secara efektif, sehingga dapat diterapkan untuk

mengenali suatu bahasa tertentu dengan lebih akurat. Proses ini melibatkan pengumpulan teks kalimat dari suatu bahasa daerah tertentu, yang kemudian akan dianalisis dan dikelompokkan berdasarkan informasi label yang diberikan. Sehingga suatu teks yang sudah dilabel dapat dikelompokkan berdasarkan label-label tersebut. Menurut (Wibawa et al., 2018) terdapat beberapa metode klasifikasi bahasa, di antaranya adalah pendekatan pembelajaran mesin seperti *Random Forest*, *Support Vector Machine* (SVM), *Naive Bayes Claasifier* (NBC) dan *K-Nearest Neighbors* (KNN). *Random Forest* merupakan perkembangan dari metode CART, yang menggabungkan teknik *bootstrap aggregating* (*bagging*) dan seleksi fitur acak. Konsep ini melibatkan pembentukan banyak pohon secara bersamaan, membentuk suatu hutan, dan kemudian melakukan analisis terhadap kumpulan pohon tersebut (Syukron & Subekti, 2018). *Support Vector Machine* (SVM) merupakan sebuah metode yang digunakan untuk melakukan prediksi, baik dalam konteks klasifikasi maupun regresi (Darmayanti et al., 2018). *Naïve Bayes Classifier* (NBC) merupakan salah satu metode yang digunakan untuk mengklasifikasikan sekumpulan teks. Algoritma ini memanfaatkan metode probabilitas dan statistik yang dikemukakan oleh ilmuwan Inggris Thomas Bayes, yaitu memprediksi probabilitas di masa depan berdasarkan pengalaman di masa sebelumnya (Setiawan et al., 2015). Algoritma *K-Nearest Neighbor* adalah metode yang termasuk dalam kategori pembelajaran terawasi (*supervised learning*). Dalam pembelajaran terawasi, fokusnya adalah untuk menemukan pola baru dalam data dengan mengaitkan pola yang sudah ada dengan data yang baru (Susandi et al., 2016).

Data yang dikumpulkan dalam penelitian ini tidak seimbang, maka dari itu data akan diseimbangkan menggunakan teknik SMOTE. *Synthetic Minority Oversampling Technique* (SMOTE) merupakan algoritma yang bekerja dengan tujuan mengontrol pemerataan distribusi data pada suatu kelas dalam dataset dengan membuat sampel buatan dari kelas minoritas (Sharfina Nabilah & Ramadhan Nur, 2022). *Confusion Matrix* merupakan teknik yang digunakan untuk mengevaluasi klasifikasi model untuk memperkirakan objek yang benar atau salah. Sebuah matriks dari prediksi akan dibandingkan dengan kelas asli yang berisi informasi aktual dan prediksi nilai klasifikasi (Pravina et al., 2019).

Penelitian mengenai klasifikasi bahasa sudah pernah dilakukan oleh (Tuhenay & Mailoa, 2021), peneliti melakukan perbandingan klasifikasi bahasa menggunakan metode NBC dan SVM dengan hasil yaitu kedua metode tersebut bagus dalam melakukan identifikasi bahasa, namun SVM lebih efektif daripada NBC. Kemudian penelitian yang dilakukan oleh (Yasir & Suraji, 2023) yaitu peneliti membandingkan akurasi beberapa metode klasifikasi, seperti *Naïve Bayes*, *Decision Tree*, dan *Random Forest*. Hasil dari penelitian tersebut menunjukkan bahwa akurasi *Naïve Bayes* mencapai 90%, *Decision Tree* 83%, dan *Random Forest* 87%. Berdasarkan penelitian-penelitian tersebut, pada penelitian ini peneliti memilih algoritma *Random Forest* dan *Support Vector Machine* (SVM) dalam menganalisis kemiripan bahasa. Kedua algoritma tersebut digunakan untuk mengidentifikasi kemiripan antara bahasa-bahasa yang berbeda dan melakukan analisis perbandingan mengenai kemiripan bahasa.

Berdasarkan masalah yang ada dalam penelitian ini, maka peneliti menggunakan beberapa bahasa berdasarkan suku yang ada di Kalimantan Barat. Beberapa bahasa yang memiliki kemiripan bahasa tersebut seperti dari rumpun dayak yaitu Bahasa Dayak Iban, Bahasa Dayak Pesaguan. Bahasa dari rumpun melayu yaitu Bahasa Melayu Pontianak, dan Bahasa Melayu Sambas. Hasil dari penelitian ini diharapkan dapat memberikan kontribusi penting dalam memahami keragaman bahasa di Indonesia, terutama dalam konteks mengenali kemiripan bahasa daerah di Kalimantan Barat.

1.2 Perumusan Masalah

Berdasarkan pada latar belakang yang telah dijelaskan, berikut beberapa rumusan masalah yang dapat menjadi fokus dalam penelitian ini:

1. Bagaimana hasil dari perbandingan antara algoritma *Random Forest* dan *Support Vector Machine*?
2. Bagaimana tingkat kemiripan antara bahasa Indonesia, bahasa Dayak Iban, bahasa Dayak Pesaguan, bahasa Melayu Pontianak, dan bahasa Melayu Sambas?

1.3 Tujuan Penelitian

Penelitian ini bertujuan mengetahui hasil perbandingan algoritma *Random Forest* dan SVM serta mengetahui tingkat kemiripan antara bahasa Indonesia, bahasa Dayak Iban, bahasa Dayak Pesaguan, bahasa Melayu Pontianak, dan bahasa Melayu Sambas.

1.4 Pembatasan Masalah

Penulis mengusulkan pembatasan masalah penelitian sebagai berikut:

1. Data yang digunakan berupa teks dari bahasa Indonesia dan beberapa bahasa daerah yaitu bahasa Melayu Pontianak, bahasa Melayu Sambas, bahasa Dayak Iban, dan bahasa Dayak Pesaguan.
2. Data yang digunakan yaitu data Non-SMOTE dan data SMOTE.
3. Algoritma yang digunakan yaitu *Random Forest* dan *Support Vector Machine* (SVM).
4. Hasil dan analisis dari penelitian ini didapatkan dari hasil *heatmap confusion matrix*.

1.5 Sistematika Penulisan

Sistematika dari penulisan skripsi ini terdiri dari lima bab yaitu, Bab I Pendahuluan, Bab II Tinjauan Pustaka, Bab III Metode Penelitian, Bab IV Hasil dan Pembahasan, dan Bab V Penutup.

BAB I PENDAHULUAN

Berisi tentang latar belakang, rumusan masalah, batasan masalah, tujuan

penelitian, metode penelitian, dan sistematika penulisan.

BAB II TINJAUAN PUSTAKA DAN LANDASAN TEORI

Berisi tentang teori-teori mengenai *Natural Language Processing (NLP)*, *Term Frequency-Inverse Document Frequency (TF-IDF)*, Bahasa Daerah, Imbalance Class, *Synthetic Minority Over-sampling Technique (SMOTE)*, algoritma *Random Forest*, algoritma *Support Vector Machine*, dan *Confusion Matrix*.

BAB III METODE PENELITIAN

Berisi tentang diagram alir dan rangkaian alur tahapan pelaksanaan penelitian, dari mulai pemahaman masalah yang diangkat oleh penulis, bagaimana persiapan *dataset*, sampai pada pembuatan model.

BAB IV HASIL DAN PEMBAHASAN

Berisi tentang hasil skenario pengujian yang sudah dilakukan. Dari hasil tersebut akan dibahas mengenai kinerja dari algoritma serta hasil kemiripan berdasarkan perhitungan algoritma *Random Forest* dan *Support Vector Machine*, beserta arah kesimpulan penelitian.

BAB V PENUTUP

Berisi tentang kesimpulan serta saran terkait penelitian yang sudah dilakukan oleh penulis.

BAB II TINJAUAN PUSTAKA

2.1 Penelitian Terkait

Dalam penelitian ini terdapat hal yang mengacu kepada beberapa penelitian yang telah dilakukan oleh peneliti-peneliti sebelumnya dengan topik terkait. Hal-hal yang mencangkup adalah klasifikasi bahasa, algoritma *Random Forest*, algoritma *Support Vector Machine* (SVM), dan *Confusion Matrix*.

(Tuhenay & Mailoa, 2021) melakukan penelitian perbandingan klasifikasi bahasa menggunakan metode *Naive Bayes Classifier* (NBC) dan *Support Vector Machine* (SVM). Peneliti bertujuan untuk mengidentifikasi bahasa dalam bentuk teks. Melalui penulisan ini akan dipaparkan hasil identifikasi bahasa Indonesia, Ambon dan Jawa dengan menggunakan teknologi komputerisasi *Machine Learning* dan memakai metode NBC dan SVM untuk menghitung nilai *accuracy* pada kedua metode tersebut agar dapat mengidentifikasi bahasa sesuai dengan teks yang dimasukkan. Penelitian ini juga melakukan perbandingan dari kedua metode tersebut agar dapat mengetahui metode apa yang lebih efektif dipakai untuk mengidentifikasi bahasa. Dari penelitian yang dilakukan, hasil yang didapatkan dengan menggunakan metode NBC dan SVM adalah keduanya bagus dalam melakukan identifikasi bahasa karena memperoleh nilai *accuracy* di atas 0,90 hanya saja melalui perhitungan *confusion matrix*, metode SVM lebih efektif dengan nilai *accuracy* 0,9634 dibandingkan dengan nilai NBC 0,9378.

Pada penelitian (Yasir & Suraji, 2023), melakukan penelitian tentang perbandingan klasifikasi *naive bayes*, *decision tree*, *random forest* terhadap analisis sentimen kenaikan biaya haji 2023 pada media sosial youtube. Penelitian ini bertujuan untuk menganalisis sentimen masyarakat dengan mengklasifikasikan komentar positif dan negatif di media sosial YouTube terkait kebijakan kenaikan biaya haji tahun 2023. Penelitian ini membandingkan akurasi beberapa metode klasifikasi, seperti *Naïve Bayes*, *Decision Tree*, dan *Random Forest*. Hasil dari penelitian ini menunjukkan bahwa akurasi *Naïve Bayes* mencapai 90%, *Decision Tree* 83%, dan *Random Forest* 87%.