



**PENERAPAN *NAÏVE BAYES* DAN *SUPPORT VECTOR MACHINE*
DALAM ANALISIS SENTIMEN PADA TEKS KESEHATAN
MENTAL DENGAN METODE PENYEIMBANGAN DATA
*OVERSAMPLING***

SKRIPSI

Untuk Memenuhi Persyaratan
Dalam Menyelesaikan Sarjana Strata-1 Ilmu Komputer

Oleh

GHISCA NUR ARIFA

NIM 2111016220022

**PROGRAM STUDI S-1 ILMU KOMPUTER
FAKULTAS MATEMATIKA DAN ILMU PENGETAHUAN ALAM
UNIVERSITAS LAMBUNG MANGKURAT
BANJARBARU**

MEI 2026

SKRIPSI

PENERAPAN *NAÏVE BAYES* DAN *SUPPORT VECTOR MACHINE* DALAM ANALISIS SENTIMEN PADA TEKS KESEHATAN MENTAL DENGAN METODE PENYEIMBANGAN DATA *OVERSAMPLING*

Oleh:

**GHISCA NUR ARIFA
2111016220022**

telah dipertahankan di depan Dosen Penguji pada tanggal 26 Mei 2026

Susunan Penguji:

Penguji

1.

Pembimbing Utama



M. Itqan Mazdadi, S.Kom, M.Kom.
NIP. 199006122019031013



Dwi Kartini, S.Kom, M.Kom.
NIP. 198704212012122003

2.

Pembimbing Pendamping



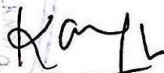
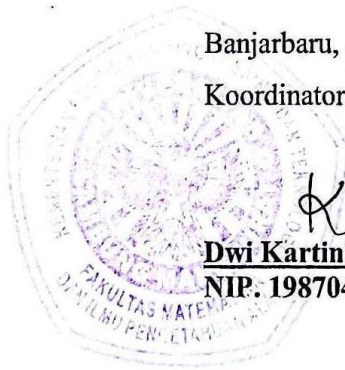
Irwan Budiman, S.T., M.Kom.
NIP. 197703252008121001



Friska Abadi, S.Kom, M.Kom.
NIP. 198809132023211010

Banjarbaru, 14 Juli 2026

Koordinator Program Studi Ilmu Komputer



Dwi Kartini, S.Kom, M.Kom.
NIP. 198704212012122003

PERNYATAAN

Dengan ini saya menyatakan bahwa dalam skripsi ini tidak terdapat karya yang pernah diajukan untuk memperoleh gelar kesarjanaan di suatu Perguruan Tinggi, dan sepanjang pengetahuan saya juga 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.

Banjarbaru, 26 Mei 2026



Ghisca Nur Arifa
NIM 2111016220022

ABSTRAK

PENERAPAN NAÏVE BAYES DAN SUPPORT VECTOR MACHINE DALAM ANALISIS SENTIMEN PADA TEKS KESEHATAN MENTAL DENGAN METODE PENYEIMBANGAN DATA OVERSAMPLING (Oleh: Ghisca Nur Arifa; Pembimbing: Muhammad Itqan Mazdadi, S.Kom, M.Kom. dan Irwan Budiman, S.T., M.Kom.; 2026; 105 halaman)

Penelitian eksperimental ini mengkaji komparasi performa antara algoritma *Support Vector Machine* (SVM) dan *Naïve Bayes* (NB) dalam mengklasifikasikan sentimen pada teks kesehatan mental yang mengalami kendala ketidakseimbangan kelas (*imbalanced data*). Dengan memanfaatkan dataset tekstual sebanyak 53.043 baris, pengujian dilakukan berbasis pembobotan TF-IDF melalui simulasi variasi rasio pembagian data (70:30, 80:20, dan 90:10) serta penerapan empat variasi teknik oversampling, meliputi *Synthetic Minority Over-sampling Technique* (SMOTE), *Borderline-SMOTE*, dan *RandomOverSampler*. Temuan eksperimen pada dataset awal mengonfirmasi superioritas SVM berkernel linear yang mencapai akurasi tertinggi sebesar 74% karena ketangguhannya dalam memetakan batas keputusan pada ruang fitur berdimensi tinggi, sementara *Naïve Bayes* hanya mencatatkan akurasi maksimal 66%. Meskipun introduksi metode *oversampling* memicu fenomena *Accuracy Paradox*—yang ditandai dengan koreksi penurunan akurasi global menjadi 67%–72% untuk SVM dan 56%–61% untuk *Naïve Bayes* akibat risiko *noise* sintetis—intervensi ini terbukti secara signifikan mendongkrak nilai sensitivitas (*recall*) dalam mendeteksi kelas minoritas yang ekstrem seperti *Personality Disorder* (2,26%). Dengan demikian, pendekatan oversampling sangat direkomendasikan apabila orientasi utama sistem ditujukan untuk mengoptimalkan akurasi identifikasi pada kategori-kategori minoritas.

Kata kunci: Analisis Sentimen, Kesehatan Mental, *Naïve Bayes*, *Support Vector Machine*, *Oversampling*.

ABSTRACT

APPLICATION OF NAÏVE BAYES AND SUPPORT VECTOR MACHINE IN SENTIMENT ANALYSIS ON MENTAL HEALTH TEXTS USING OVERSAMPLING DATA BALANCING METHOD (By: Ghisca Nur Arifa; Supervisors: Muhammad Itqan Mazdadi, S.Kom, M.Kom. and Irwan Budiman, S.T., M.Kom.; 2026; 105 pages)

This experimental study investigates the performance comparison between Support Vector Machine (SVM) and Naïve Bayes (NB) algorithms in classifying sentiments from mental health text data characterized by severe class imbalance. Utilizing a textual dataset of 53,043 rows, the evaluation was conducted using TF-IDF weighting across three data-splitting ratios (70:30, 80:20, and 90:10) combined with four oversampling techniques, including Synthetic Minority Over-sampling Technique (SMOTE), Borderline-SMOTE, and RandomOverSampler. Experimental findings on the baseline dataset confirm the superiority of the linear-kernel SVM, which achieved a peak accuracy of 74% due to its robustness in defining decision boundaries within high-dimensional feature spaces, whereas Naïve Bayes only reached a maximum accuracy of 66%. Although the introduction of oversampling methodologies triggered the Accuracy Paradox phenomenon—characterized by a slight decline in global accuracy to 67%–72% for SVM and 56%–61% for Naïve Bayes caused by synthetic noise—this intervention significantly boosted the sensitivity (recall) in detecting extreme minority classes such as Personality Disorder (2.26%). Consequently, the oversampling approach is highly recommended when the system's primary orientation is aimed at optimizing identification accuracy within minority categories.

Keywords: Sentiment Analysis, Mental Health, Social Media, Naïve Bayes, Support Vector Machine, Oversampling

PRAKATA

Assalamu'alaikum Warrahmatullahi Wabarakatuh.

Puji syukur penulis panjatkan ke hadirat Allah SWT karena atas berkat rahmat dan hidayah-Nya penulis dapat menyelesaikan skripsi yang berjudul “Penerapan *Naïve Bayes* dan *Support Vector Machine* dalam Analisis Sentimen pada Teks Kesehatan Mental dengan Metode Penyeimbangan Data *Oversampling*” untuk memenuhi syarat dalam menyelesaikan pendidikan program S-1 Ilmu Komputer, Fakultas Matematika dan Ilmu Pengetahuan Alam, Universitas Lambung Mangkurat. Tak lupa pula penulis panjatkan shalawat serta salam pada Nabi Muhammad SAW beserta para sahabat, keluarga dan pengikut beliau hingga *yaumul qiyamah*.

Pada lembar ini penulis ingin mengucapkan ucapan terima kasih kepada pihak-pihak yang sangat mendukung penulis dalam pembuatan dan penyusunan skripsi ini, adapun yang dimaksud adalah sebagai berikut:

1. Keluarga yang selalu mendukung dalam proses penyelesaian skripsi dan menguatkan batin di saat kehilangan arah, yaitu Alm. Ayah, Ibu, Kakak Difo, Kakak Fadel, Kakak Novi dan Affan.
2. Bapak Muhammad Itqan Mazdadi, S.Kom, M.Kom. dan Bapak Irwan Budiman, S.T., M.Kom. selaku dosen pembimbing utama dan pendamping yang turut serta membantu dan meluangkan waktu demi kelancaran dalam penyelesaian skripsi ini.
3. Ibu Dwi Kartini, S.Kom., M.Kom. selaku Ketua Program Studi Ilmu Komputer FMIPA ULM, dan Bapak Triando Hamonangan Saragih, S.Kom, M.Kom. selaku dosen pembimbing akademik yang telah memberikan izin sehingga pengerjaan skripsi ini dapat terselesaikan.
4. Seluruh Dosen dan Staff Program Studi Ilmu Komputer FMIPA ULM atas ilmu dan bantuan yang diberikan selama ini yang sangat bermanfaat.
5. Teman-teman dan sahabat-sahabat keluarga Ilmu Komputer angkatan 2021 yang selalu memberikan dukungan dan mengingatkan serta mendoakan dalam proses mengerjakan skripsi

6. Ucapan terima kasih kepada Kak Felia, Sethos, Nari, Hayi, Collei, Onti Furina, Tao, Mico, Ayaka, Scar, Mbaknyu, Retatang, Kak Sho dan Lili yang sudah menemani penulis dalam suka serta duka di masa pengerjaan skripsi
7. Semua pihak yang tidak dapat disebutkan satu persatu yang telah turut membantu dalam penyelesaian skripsi ini.

Akhir kata penulis menyadari sepenuhnya bahwa penulisan ini jauh dari sempurna, namun penulis mengharapkan bantuan serupa berupa saran dan kritik yang membangun dari semua pihak demi kesempurnaan dan mutu penulisan skripsi ini. Semoga tulisan ini dapat bermanfaat bagi ilmu pengetahuan dan pembaca khususnya serta mendapat keridhaan Allah SWT.

Banjarbaru, 26 Mei 2026



Ghisca Nur Arifa

DAFTAR ISI

HALAMAN JUDUL.....	i
HALAMAN PENGESAHAN.....	Error! Bookmark not defined.
PERNYATAAN.....	iii
ABSTRAK	iv
<i>ABSTRACT</i>	v
PRAKATA.....	vi
DAFTAR ISI.....	viii
DAFTAR TABEL	xi
DAFTAR GAMBAR	xiii
DAFTAR LAMPIRAN	xiv
BAB I PENDAHULUAN.....	1
1.1 Latar Belakang.....	1
1.2 Rumusan Masalah.....	4
1.3 Tujuan Penelitian	4
1.4 Manfaat Penelitian	4
1.5 Batasan Masalah	5
1.6 Sistematika Penulisan	5
BAB II TINJAUAN PUSTAKA.....	7
2.1 Kajian Terdahulu	7
2.2 Landasan Teori	12
2.2.1 Media Sosial	12
2.2.2 Analisis Sentimen	13
2.2.3 Preprocessing	13

2.2.4	Term Frequency-Inverse Document Frequency (TF-IDF)	14
2.2.5	Algoritma <i>Naïve Bayes</i>	16
2.2.6	Algoritma <i>Support Vector Machine</i>	17
2.2.7	Synthetic Minority Over-sampling Technique (SMOTE)	19
2.2.8	Random Oversampling	21
2.2.9	Borderline-SMOTE	22
2.2.10	Evaluasi Kinerja Model	24
BAB III METODE PENELITIAN		26
3.1	Alat Penelitian.....	26
3.2	Bahan Penelitian	26
3.2	Prosedur Penelitian	27
BAB IV HASIL DAN PEMBAHASAN.....		30
4.1	Hasil	30
4.1.1	Pengumpulan Data	30
4.1.2	Analisis Pola Kata pada Setiap Label Sentimen	32
4.1.3	Pre-processing Data	34
4.1.4	Ekstraksi Fitur	35
4.1.5	Pembagian Data	36
4.1.6	Penyeimbangan Data (<i>Oversampling</i>)	37
4.1.7	Pelatihan Model Klasifikasi.....	38
4.1.8	Analisis dalam Algoritma <i>Naïve Bayes</i>	40
4.1.9	Analisis dalam Algoritma <i>Support Vector Machine</i>	46
4.2	Pembahasan	55
BAB V PENUTUP		58
5.1	Kesimpulan	58

5.2	Saran	59
	DAFTAR PUSTAKA	60
	LAMPIRAN	65

DAFTAR TABEL

Tabel 1. Keaslian Penelitian.....	9
Tabel 2. Perancangan Penelitian	12
Tabel 3. Confusion Matrix	24
Tabel 4. Data teks kesehatan mental	30
Tabel 5. Distribusi Label Unik.....	31
Tabel 6. Tahapan pre-processing data yang dilakukan	34
Tabel 7. Data teks kesehatan mental yang telah dilakukan <i>pre-processing</i>	34
Tabel 8. Hasil TF-IDF dari nilai kata teringgi	35
Tabel 9. Pembagian data <i>train</i> dan <i>testing</i> pada skenario 70:30.....	36
Tabel 10. Pembagian data <i>train</i> dan <i>testing</i> pada skenario 80:20.....	37
Tabel 11. Pembagian data <i>train</i> dan <i>testing</i> pada skenario 90:10.....	37
Tabel 12. Jumlah data pada tiap label yang akan dilatih.....	38
Tabel 13. Tokenisasi pada salah satu data dan frekuensi kata dalam data.....	39
Tabel 14. Nilai TF-IDF dan indeks fitur tiap kata dari sampel data	39
Tabel 15. Nilai Probabilitas Prior dan <i>Class Log Prior</i> pada tiap label.....	41
Tabel 16. Nilai <i>Log</i> , Probabilitas <i>Likelihood</i> dan kontribusi kata dalam sampel .	42
Tabel 17. Nilai <i>Log Prior</i> , Total Kontribusi Kata dan <i>Log Score Manual</i>	43
Tabel 18. Nilai eksponensial dari setiap label.....	44
Tabel 19. Nilai probabilitas prediksi dan presentase dari setiap label	46
Tabel 20. Nomor urut index label beserta label	47
Tabel 21. Nilai TF-IDF, bobot SVM dan kontribusi kata dari salah satu sampel.	48
Tabel 22. Nilai <i>intercept</i> , kontribusi dan <i>decision score manual</i> setiap label	50
Tabel 23. Nilai <i>decision score</i> , <i>softmax-like</i> dan persentase setiap label.....	52

Tabel 24. Evaluasi Model Tanpa <i>Oversampling</i>	53
Tabel 25. Evaluasi Model dengan <i>Oversampling SMOTE</i>	53
Tabel 26. Evaluasi Model dengan <i>Oversampling Borderline-SMOTE</i>	54
Tabel 27. Evaluasi Model dengan <i>Oversampling RandomOverSampler</i>	54
Tabel 28. Perbandingan antara NB dan SVM diberbagai rasio dan skenario	55

DAFTAR GAMBAR

Gambar 1. Alur Penelitian.....	27
Gambar 2. Visualisasi Label Unik	31
Gambar 3. Perbandingan nilai akurasi NB dan SVM	56

DAFTAR LAMPIRAN

Lampiran 1. Dataset

Lampiran 2. *Install Library* Pihak Ketiga

Lampiran 3. Import Modul dan *Install Library* Utama

Lampiran 4. Pengunduhan *Resource NLTK* dan Inisialisasi Komponen

Lampiran 5. Pembuatan Fungsi *Pipeline Preprocessing Text*

Lampiran 6. Memuat *Dataset*, Eksekusi Fungsi dan Ekspor Hasil *Cleaned Data*

Lampiran 7. Program yang memproses tanpa Oversampling

Lampiran 8. Program yang memproses SMOTE

Lampiran 9. Program yang memproses RandomOverSampler

Lampiran 9. Output Pre-processing

Lampiran 10. Output Top Words TF-IDF Each Status (Anxiety)

Lampiran 11. *Output Top Words TF-IDF Each Status (Normal & Depression)*

Lampiran 12. *Output Top Words TF-IDF Each Status (Suicidal & Stress)*

Lampiran 13. *Output Top Words TF-IDF Each Status (Bipolar & Personality D.)*