



KOMPARASI ALGORITMA *NAÏVE BAYES* DAN *K-NEAREST NEIGHBOR (K-NN)* UNTUK KLASIFIKASI *FAKE JOB ADVERTISEMENT* DI *PLATFORM MEDIA SOSIAL*

SKRIPSI

**Untuk memenuhi persyaratan
dalam menyelesaikan program sarjana Strata-1 Matematika**

**Oleh:
SITI HARDIYANTI
NIM. 2211011120009**

**JURUSAN MATEMATIKA
FAKULTAS MATEMATIKA DAN ILMU PENGETAHUAN ALAM
UNIVERSITAS LAMBUNG MANGKURAT
BANJARBARU
2026**

HALAMAN PENGESAHAN

SKRIPSI

**KOMPARASI ALGORITMA *NAÏVE BAYES* DAN *K-NEAREST NEIGHBOR*
(*K-NN*) UNTUK KLASIFIKASI *FAKE JOB ADVERTISEMENT*
DI PLATFORM MEDIA SOSIAL**

Oleh:

Siti Hardiyanti

NIM.2211011120009

telah dipertahankan di depan Dosen Penguji pada tanggal 10 Juli 2026.

Susunan Dosen Penguji:

Pembimbing I



Oni Soesanto, S.Si., M.Si.

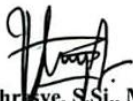
NIP. 197301262005011003

Dosen Penguji

1. Saman Abdurrahman, S.Si., M.Sc. (✓²)

2. Dr. Muhammad Ahsar Karim, S.Si., M.Sc. (✓)

Pembimbing II



Thresye, S.Si., M.Si.

NIP. 197205042000122002

Banjarbaru, 23 Juli 2026

Ketua Jurusan Matematika

FMIPA ULM

Dr. Na'imah Hijriati, S.Si., M.Si.

NIP. 197911222008012013



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, 23 Juli 2026



Siti Hardiyanti

2211011120009

ABSTRAK

KOMPARASI ALGORITMA *NAÏVE BAYES* DAN *K-NEAREST NEIGHBOR (K-NN)* UNTUK KLASIFIKASI *FAKE JOB ADVERTISEMENT* DI *PLATFOR MEDIA SOSIAL* (oleh : Siti Hardiyanti ; Oni Soesanto, Theresye, 100 halaman)

Maraknya penyebaran iklan lowongan pekerjaan melalui media sosial meningkatkan kebutuhan akan metode yang mampu membedakan iklan asli dan *Fake Job Advertisement* secara otomatis. Penelitian ini bertujuan merepresentasikan teks *Fake Job Advertisement* ke dalam bentuk numerik melalui tahapan *text pre-processing* dan pembobotan *Term Frequency-Inverse Document Frequency* (TF-IDF) sebagai dasar ekstraksi fitur, serta membandingkan efektivitas algoritma *Bernoulli Naïve Bayes*, *Multinomial Naïve Bayes*, dan *K-Nearest Neighbor (K-NN)* dalam mengklasifikasikan *Fake Job Advertisement* pada platform media sosial. Dataset yang digunakan terdiri atas 246 data yang diperoleh dari berbagai platform media sosial dan sumber sekunder, kemudian dilabeli ke dalam dua kelas, yaitu penipuan dan asli. Tahapan *text pre-processing* meliputi *data cleaning*, *case folding*, *tokenization*, *normalization*, *stopword removal*, dan *stemming*. Selanjutnya, data diekstraksi menggunakan TF-IDF dan dibagi ke dalam tiga skenario pelatihan dan pengujian, yaitu 90%:10%, 80%:20%, dan 70%:30%. Evaluasi dilakukan menggunakan *confusion matrix* berdasarkan metrik *accuracy*, *precision*, *recall*, dan *F1-score*. Hasil penelitian menunjukkan bahwa algoritma *K-Nearest Neighbor (K-NN)* memberikan performa terbaik dibandingkan *Bernoulli Naïve Bayes* dan *Multinomial Naïve Bayes*. Performa terbaik diperoleh pada skenario 90% data latih dan 10% data uji dengan *accuracy* sebesar 92,00%, serta *precision*, *recall*, dan *F1-score* masing-masing sebesar 92,31%, sehingga *K-NN* merupakan algoritma yang paling efektif dalam mengklasifikasikan *Fake Job Advertisement* pada platform media sosial

Kata kunci: *Fake Job Advertisement*, *Text Pre-processing*, *TF-IDF*, *Naïve Bayes*, *K-Nearest Neighbor (K-NN)*.

ABSTRACT

COMPARISON OF NAÏVE BAYES AND K-NEAREST NEIGHBOR (K-NN) ALGORITHMS FOR FAKE JOB ADVERTISEMENT CLASSIFICATION ON SOCIAL MEDIA PLATFORMS (by: Siti Hardiyanti; Oni Soesanto, Theresye, 100 pages)

The widespread dissemination of job vacancy advertisements through social media has increased the need for methods capable of automatically distinguishing legitimate job advertisements from Fake Job Advertisements. This study aims to represent Fake Job Advertisement texts into numerical representations through text pre-processing and Term Frequency–Inverse Document Frequency (TF-IDF) weighting as the basis for feature extraction, as well as to compare the effectiveness of the Bernoulli Naïve Bayes, Multinomial Naïve Bayes, and K-Nearest Neighbor (K-NN) algorithms in classifying Fake Job Advertisements on social media platforms. The dataset consisted of 246 job advertisement records collected from various social media platforms and secondary sources, which were manually labeled into two classes: fraudulent and legitimate. The text pre-processing stages included data cleaning, case folding, tokenization, normalization, stopword removal, and stemming. The data were then transformed using TF-IDF and divided into three training and testing scenarios: 90%:10%, 80%:20%, and 70%:30%. Model performance was evaluated using a confusion matrix based on accuracy, precision, recall, and F1-score. The results showed that the K-Nearest Neighbor (K-NN) algorithm outperformed the Bernoulli Naïve Bayes and Multinomial Naïve Bayes algorithms across all data-splitting scenarios. The best performance was achieved under the 90% training and 10% testing scenario, with an accuracy of 92.00% and precision, recall, and F1-score of 92.31%, indicating that K-NN is the most effective algorithm for classifying Fake Job Advertisements on social media platforms.

Keywords: *Fake Job Advertisement, Text Pre-processing, TF-IDF, Naïve Bayes, K-Nearest Neighbor (K-NN).*

PRAKATA

Puji syukur penulis panjatkan ke hadirat Allah Subhānahu Wa Ta‘ālā atas limpahan rahmat, taufik, dan hidayah-Nya sehingga penulis dapat menyelesaikan skripsi yang berjudul **“Komparasi Algoritma Naïve Bayes dan K-Nearest Neighbor (K-NN) untuk Klasifikasi Fake Job Advertisement di Platform Media Sosial.**

Shalawat serta salam semoga senantiasa tercurah kepada Nabi Muhammad ﷺ ‘Alaihi Wa Sallam, beserta keluarga, sahabat, dan para pengikut beliau hingga akhir zaman.

Skripsi ini disusun untuk memenuhi salah satu persyaratan dalam menyelesaikan Program Sarjana (S-1) pada Program Studi Matematika, Fakultas Matematika dan Ilmu Pengetahuan Alam, Universitas Lambung Mangkurat. Penulis menyadari bahwa penyusunan skripsi ini dapat terselesaikan berkat bantuan, arahan, serta dukungan dari berbagai pihak. Oleh karena itu, penulis menyampaikan rasa hormat dan terima kasih kepada:

1. Bapak Prof. Ir. Sunardi, S.Si., M.Sc., Ph.D., selaku Dekan Fakultas Matematika dan Ilmu Pengetahuan Alam Universitas Lambung Mangkurat.
2. Ibu Dr. Na’imah Hijriati, S.Si., M.Si., selaku Koordinator Program Studi Matematika beserta seluruh dosen dan staf Program Studi Matematika.
3. Bapak Nurul Huda S.Si., M.Si. selaku Dosen Penasehat Akademik yang telah memberikan arahan dan bimbingan selama menjalani perkuliahan.
4. Bapak Oni Soesanto, S.Si., M.Si. selaku Dosen Pembimbing I dan Ibu Thresye, S.Si., M.Si. selaku Dosen Pembimbing II, yang telah memberikan bimbingan, nasihat, motivasi, dan petunjuk dalam pelaksanaan penelitian serta penyusunan skripsi ini.
5. Bapak Saman Abdurrahman, S.Si., M.Sc. selaku Dosen Penguji I dan Bapak Dr. Muhammad Ahsar Karim, S.Si., M.Sc. selaku Dosen Penguji II yang telah memberikan kritik dan saran demi kesempurnaan skripsi ini.

6. Kedua orang tua serta keluarga besar yang selalu memberikan doa, dukungan, dan motivasi selama proses penyusunan skripsi.
7. Teman-teman dan seluruh pihak yang tidak dapat penulis sebutkan satu per satu, yang telah memberikan bantuan dan dukungan dalam proses penelitian dan penulisan skripsi ini.

Penulis menyadari bahwa skripsi ini masih memiliki kekurangan. Oleh karena itu, kritik dan saran yang membangun sangat penulis harapkan demi perbaikan di masa mendatang . Akhir kata, penulis berharap skripsi ini dapat memberikan manfaat bagi para pembaca dan semua pihak yang berkepentingan.

Banjarbaru, 23 Juli 2026



Siti Hardiyanti

2211011120009

DAFTAR ISI

	Halaman
ABSTRAK.....	i
ABSTRACT	ii
PRAKATA	iii
DAFTAR ISI	v
DAFTAR TABEL	vii
DAFTAR GAMBAR	viii
DAFTAR LAMPIRAN	ix
ARTI LAMBANG DAN SINGKATAN	x
BAB I PENDAHULUAN	1
1.1. Latar Belakang	1
1.2. Tujuan Penelitian	3
1.3. Sistematika Penulisan	4
BAB II TINJAUAN PUSTAKA	2
2.1 Platform Digital.....	2
2.2 Iklan Lowongan Pekerjaan.....	6
2.3 Data Mining	7
2.4 Klasifikasi	8
2.5 Text Pre-Processing	9
2.6 <i>Naïve Bayes Classifier</i>	11
2.6.1 <i>Bernoulli Naive Bayes</i>	13
2.6.2 <i>Multinomial Naïve Bayes</i>	15
2.7 <i>K-Nearest Neighbor (K-NN)</i>	17
2.8 <i>Term Frequency-Inverse Document Frequency (TF-IDF)</i>	18
2.9 <i>Confusion Matrix</i>	19
BAB III METODE PENELITIAN.....	22
3.1 Sumber Data.....	22
3.2 Prosedur Penelitian	22
3.3 Flowchart Penelitian	26
BAB IV HASIL DAN PEMBAHASAN	27

4.1	Penerapan Text Pre-Processing dan Pembobotan TF – IDF	27
4.1.1	Pengumpulan Data	27
4.1.2	Pelabelan Data	28
4.1.3	<i>Text Pre Processing</i>	31
4.1.3.1	<i>Data Cleaning</i>	31
4.1.3.2	<i>Case Folding</i>	32
4.1.3.3	<i>Tokenization</i>	32
4.1.3.4	<i>Normalization</i>	33
4.1.3.5	<i>Stopword Removal</i>	34
4.1.3.6	<i>Stemming</i>	35
4.1.4	<i>Word Cloud</i>	36
4.1.4.1	<i>Word Cloud Penipuan</i>	36
4.1.4.2	<i>Word Cloud Asli</i>	37
4.1.5	Pembagian Data Training dan Testing	38
4.1.6	Pembobotan (<i>TF IDF</i>)	39
4.2	Perbandingan Naïve Bayes dan K-Nearest Neighbor (K-NN)	44
4.2.1	Proses Perhitungan Klasifikasi Bernaulli Naïve Bayes	44
4.2.2	Proses Perhitungan Klasifikasi Multinomial Naïve Bayes	47
4.2.3	Proses Perhitungan Klasifikasi K-Nearest Neighbor K-NN	50
4.2.4	Implementasi Klasifikasi Bernaulli Naïve Bayes	53
4.2.5	Implementasi Klasifikasi Multinomial Naïve Bayes	54
4.2.6	Implementasi Klasifikasi K-Nearest Neighbor (KNN)	56
4.2.7	Evaluasi Hasil Bernaulli Naïve Bayes	57
4.2.8	Evaluasi Hasil Multinomial Naïve Bayes	59
4.2.9	Evaluasi Hasil K-Nearest Neighbor (K-NN)	60
4.2.10	Interpretasi Hasil Perbandingan	62
BAB V	PENUTUP	65
5.1	Kesimpulan	65
5.2	Saran	66
DAFTAR PUSTAKA		67
LAMPIRAN		71

DAFTAR TABEL

	Halaman
Tabel 2.1 <i>Confusion matrix</i>	19
Tabel 4.1. Dataset yang akan digunakan	27
Tabel 4.2. Dataset yang sudah diberi lebel.....	29
Tabel 4.3 Hasil Proses data <i>Cleaning</i>	31
Tabel 4.4 Hasil Proses <i>Case Folding</i>	32
Tabel 4.5 Hasil Proses <i>Tokenization</i>	33
Tabel 4.6 Hasil Proses <i>Normalization</i>	34
Tabel 4.7 Hasil Proses <i>Stopword removal</i>	35
Tabel 4.8 Hasil Proses <i>Stemming</i>	36
Tabel 4.9 Pembagian Data <i>Tranning</i> dan <i>Testing</i>	39
Tabel 4.10 Data Kedua belas Dalam Data <i>Training</i> Skenario Pertama	40
Tabel 4.11 Hasil <i>TF-IDF</i> Data Kedua belas.....	41
Tabel 4.12 Pembagian Dataset	44
Tabel 4.13 Probabilitas setiap <i>term</i> pada setiap kelas	46
Tabel 4.14 Probabilitas setiap <i>term</i> pada setiap kelas	49
Tabel 4.15 Nilai <i>TF-IDF</i> masing – masing dataset	51
Tabel 4.16 Hasil Perkalian skalar <i>TF-IDF</i> data latih dan Data uji.....	51
Tabel 4.17 Total kuadrat dan akar kuadrat <i>TF -IDF</i> data latih dan data uji.....	52
Tabel 4.18 Hasil <i>Cosine Simalirity</i>	53
Tabel 4.19 Hasil jarak data uji.....	53
Tabel 4.20 Hasil Klasifikasi <i>Bernaulli Naïve Bayes</i>	54
Tabel 4.21 Hasil Klasifikasi <i>Multinomial Naïve Bayes</i>	55
Tabel 4.22 Hasil Klasifikasi <i>K-Nearest Neighbor (KNN)</i>	56
Tabel 4.23 Hasil Evaluasi.....	62

DAFTAR GAMBAR

	Halaman
Gambar 3.1 Flowchart Penelitian	26
Gambar 4.1 Jumlah Data Lebel Penipuan dan Asli.....	30
Gambar 4.2 <i>Word Cloud</i> Lebel Penipuan	37
Gambar 4.3 <i>Word Cloud</i> Lebel Asli.....	38

DAFTAR LAMPIRAN

- Lampiran 1.** Tabel Data Penelitian dan Lebel
- Lampiran 2.** Data Hasil Tahapan *Text Pre Processing*
- Lampiran 3.** Pembobotan Tf IDF untuk Skenario 1 (90 % : 10%)
- Lampiran 4.** Pembobotan Tf IDF untuk Skenario 1 (80 % : 20%)
- Lampiran 5.** Pembobotan Tf IDF untuk Skenario 1 (70 % : 30%)
- Lampiran 6.** *Syntax colab Install Library, Import Data, dan Text Pre-processing*
- Lampiran 7.** *Syntax Colab Word Cloud*
- Lampiran 8.** *Syntax Colab Splitting Data*
- Lampiran 9.** *Syntax Colab TF IDF*
- Lampiran 10.** *Syntax Colab Bernaulli Naive Bayes*
- Lampiran 11.** *Syntax Colab Multinomial Naive Bayes*
- Lampiran 12.** *Syntax Colab K-Nearest Neighbor (K-NN)*

ARTI LAMBANG DAN SINGKATAN

c	Variabel target atau kelas
d	Data atau dokumen dengan kelas yang belum diketahui
$P(c d)$	Peluang kelas c dengan d diketahui (<i>posterior Probability</i>)
$P(d c)$	Peluang data d dengan kelas c diketahui, (<i>likelihood</i>)
$P(c)$	Peluang awal kelas c (<i>prior probability</i>)
$P(d)$	Peluang data d (<i>evidence</i>)
t	kata (<i>Term</i>)
t_k	<i>Term</i> ke- k
V	Himpunan seluruh kata (<i>vocabulary</i>) dalam data latih
n_d	Jumlah total kata dalam data d
e_i	Fitur ke- i
M	Jumlah total fitur
$P(t_k c)$	Peluang kemunculan (<i>term</i>) t_k pada data dengan kelas c
N_c	Jumlah data latih dalam kelas c
N	Jumlah total data latih
T_{ct}	Jumlah kemunculan kata (<i>term</i>) t pada data latih dengan kelas c
B	Jumlah kata (<i>term</i>) unik pada seluruh kelas
N_{ct}	Jumlah data latih dalam kelas c yang memuat <i>term</i> t
$W_{t,d}$	Bobot <i>term</i> (t) terhadap data (d)
$tf_{t,d}$	Frekuensi kemunculan <i>term</i> t dalam data d
D	Jumlah total dokumen
df_t	Jumlah dokumen yang mengandung <i>term</i> t
$\propto$	Proporsional
KNN	<i>K-Nearest Neighbor</i>
TP	<i>True Positive</i> , Data positif yang diprediksi benar
TN	<i>True Negative</i> , Data negatif yang diprediksi benar
FP	<i>False Positive</i> , Data negatif yang diprediksi sebagai positif
FN	<i>False Negative</i> , Data positif yang diprediksi sebagai negati