

Klasifikasi Multilabel pada Topik ayat Al-Qur'an Menggunakan *Random Forest* dan *Naïve Bayes*

1st Imran Zulkarnaen
Fakultas Informatika
Universitas Telkom
Bandung, Indonesia

imranzulkarnaen@student.telkomuniversity.ac.id

2nd Kemas Muslim Lhaksmana
Fakultas Informatika
Universitas Telkom
Bandung, Indonesia

kemasmuslim@telkomuniversity.ac.id

Abstrak — Al-Qur'an, sebagai kitab suci umat Islam, menyimpan makna yang mendalam, mencakup aspek akidah, ibadah, dan etika sosial. Namun, kerumitan bahasa dalam Al-Qur'an menimbulkan tantangan dalam pengelompokan ayat-ayatnya ke dalam kategori tematik tertentu, terutama dengan pendekatan tradisional yang sering kali tidak dapat menggali hubungan semantik antar kata secara mendalam. Untuk mengatasi tantangan ini, penelitian ini mengembangkan sistem klasifikasi multilabel yang berbasis graph mining, dengan memanfaatkan pengukuran centrality. Sistem tersebut melibatkan pembuatan graf kata untuk merepresentasikan hubungan antar kata, serta penerapan algoritma random forest dan naïve bayes dalam mengklasifikasikan ayat-ayat Al-Qur'an ke dalam delapan kategori tematik. Proses pengolahan data mencakup penghapusan kata henti (stopwords), tokenisasi, dan ekstraksi fitur berdasarkan centrality, seperti closeness, betweenness, dan eigenvector. Hasil penelitian menunjukkan bahwa penggunaan betweenness centrality dengan penggunaan kata henti memberikan performa terbaik, dengan nilai Hamming loss sebesar 0.1631 pada random forest. Temuan ini menekankan keunggulan pendekatan berbasis graf dalam memahami hubungan kompleks antar kata dalam teks Al-Qur'an serta berkontribusi pada pengembangan metode klasifikasi tematik berbasis teknologi yang lebih efisien.

Kata kunci— klasifikasi Multilabel, Tematik, Al-Qur'an, Graf, Sentralitas, Graph Mining, Hamming Loss

I. PENDAHULUAN

Al-Qur'an, sebagai kitab suci umat Islam, memiliki kedalaman makna yang mencakup berbagai aspek kehidupan, meliputi akidah, ibadah, serta etika sosial [1]. Pemahaman terhadap ayat-ayat Al-Qur'an secara tematik menjadi sangat penting untuk membantu individu dan masyarakat dalam memahami pesan-pesan yang terkandung di dalamnya secara lebih terstruktur. Dalam era digital yang semakin maju, teknologi informasi menyediakan berbagai metode untuk mendukung studi Al-Qur'an, salah satunya melalui teknik pengklasifikasian teks.

Klasifikasi teks merupakan sebuah bidang ilmu yang berkembang pesat dalam era data besar (big data) [2]. Metode ini memungkinkan pengelompokan informasi secara otomatis berdasarkan pola-pola tertentu, yang selanjutnya mempermudah analisis terhadap volume data yang besar. Namun, klasifikasi teks dalam konteks keagamaan, seperti pada ayat-ayat Al-Qur'an, menghadapi tantangan yang unik

[3]. Bahasa Al-Qur'an yang kompleks, kaya akan makna kontekstual dan hubungan semantik, memerlukan pendekatan yang lebih canggih daripada sekadar analisis statistik tradisional.

Perkembangan metode berbasis kecerdasan buatan, seperti machine learning dan graph mining, membuka peluang baru dalam pengelolaan teks Al-Qur'an. Dengan memanfaatkan representasi graf untuk menganalisis hubungan antar kata dan algoritma prediksi yang canggih, diharapkan proses klasifikasi ayat-ayat Al-Qur'an dapat dilakukan dengan tingkat akurasi yang lebih tinggi [4]. Penelitian ini berupaya untuk mengeksplorasi metode tersebut, tidak hanya sebagai kontribusi akademik, tetapi juga sebagai solusi praktis dalam mendukung studi tematik Al-Qur'an di era modern.

Klasifikasi ayat-ayat Al-Qur'an ke dalam kategori tematik merupakan salah satu permasalahan utama dalam bidang analisis teks keagamaan. Ayat-ayat Al-Qur'an memiliki karakteristik bahasa yang kompleks, kaya akan makna kontekstual, serta mengandung hubungan semantik yang erat antara kata-kata di dalamnya [5]. Dalam upaya untuk meningkatkan pemahaman terhadap Al-Qur'an secara lebih terstruktur, klasifikasi ini dapat memberikan manfaat yang signifikan, terutama dalam mempermudah studi tematik bagi para peneliti, akademisi, dan pembaca umum.

Pendekatan yang diadopsi dalam penelitian ini, yaitu penambangan graf (*graph mining*), menawarkan metode inovatif untuk menganalisis relasi antar kata dalam bentuk graf. Graf tersebut memfasilitasi pemetaan hubungan semantik melalui metrik sentralitas, seperti *degree*, *betweenness*, dan *closeness*, yang berperan penting dalam identifikasi kata kunci dan struktur yang signifikan dalam teks. Berdasarkan penelitian sebelumnya, penerapan teknik penambangan graf dalam analisis teks dapat memperdalam pemahaman terhadap struktur semantik yang kompleks [6]. Teknik ini menjadi menarik karena memiliki kemampuan untuk merepresentasikan hubungan antar kata secara visual dan analitis, sehingga memberikan perspektif baru dalam memahami teks suci yang kompleks, seperti Al-Qur'an. Selain itu, penelitian menunjukkan bahwa visualisasi tematik berbasis *knowledge graph* dapat memberikan kontribusi dalam pengambilan sumber daya ilmiah dan analisis yang lebih efisien [7].

Penggunaan algoritma *Random Forest* menawarkan fleksibilitas dan stabilitas dalam klasifikasi data berlabel ganda, dengan hasil yang menunjukkan akurasi mencapai

90,77% dalam analisis sentimen [8]. Sementara itu, *Naïve Bayes* merupakan algoritma yang sederhana namun efektif untuk pengolahan teks, yang dapat berfungsi sebagai bahan perbandingan. Penelitian sebelumnya juga menunjukkan bahwa *Naïve Bayes* sering kali lebih unggul dalam klasifikasi data teks jika dibandingkan dengan algoritma lain, seperti *K-Nearest Neighbor* [9]. Evaluasi kinerja ketiga algoritma menggunakan metrik *Hamming loss* dirancang untuk memberikan wawasan yang komprehensif tentang efektivitas pendekatan yang diusulkan. Hasil perbandingan ini dapat memberikan informasi penting terkait pemilihan algoritma yang tepat untuk tugas klasifikasi tertentu.

II. KAJIAN TEORI

Tabel 1 merupakan tinjauan pustaka berdasarkan hasil studi penelitian sebelumnya. Daftar referensi berikut diharapkan mampu menjadi panduan dan berkontribusi dalam kesuksesan penelitian ini.

TABEL 1
Literature Review

Nomor Referensi	Dataset	Metode	Hasil Kinerja
[10]	Menggunakan data Juz 30 dari Al-Qur'an.	<i>Decision Tree</i> , <i>Support Vector Machine</i> (SVM) dan <i>Naïve Bayes</i> .	Hasil akurasi dari masing-masing algoritma, yaitu <i>Decision Tree</i> dengan akurasi 74,34% , SVM dengan akurasi 75,84% , dan <i>Naïve Bayes</i> dengan akurasi 66,29% .
[11]	Dataset yang dianalisis mencakup 6.236 ayat Al-Qur'an dalam bahasa Arab.	CNN (<i>Convolutional Neural Network</i>) dan RNN (<i>Recurrent Neural Network</i>)	Tingkat akurasi yang diperoleh adalah CNN dengan <i>hamming loss</i> 0.1615 dan RNN dengan <i>Hamming loss</i> sebesar 0.1694 .
[12]	Data yang dianalisis merupakan teks Al-Qur'an berbahasa Arab yang diperoleh dari Tanzil.net.	<i>Naïve Bayes</i> dan <i>Support Vector Machine</i> (SVM)	<i>Naïve Bayes</i> mempunyai <i>hamming loss</i> sebesar 0.2077 , dan SVM memiliki <i>hamming loss</i> sebesar 0.1540 .
[13]	Data terdiri dari Surah al-Baqarah Al-Qur'an.	<i>Naïve Bayes</i> , <i>J48</i> , <i>K-Nearest Neighbor</i> (KNN) dan <i>Support Vector Machine</i> (SVM)	Dalam penelitian ini, SVM dan <i>J48</i> memiliki akurasi diatas 85% , KNN dengan tingkat akurasi 71% dan <i>Naïve Bayes</i> mempunyai akurasi dibawah SVM dan <i>J48</i> .
[14]	Kumpulan data yang digunakan dalam penelitian ini merupakan terjemahan Al-Qur'an oleh Abdullah Yusuf Ali, yang mencakup 6.236 ayat.	<i>Support Vector Machine</i> (SVM), <i>Naïve Bayes</i> , <i>K-Nearest Neighbor</i> (KNN), dan <i>Decision Tree</i> .	SVM, <i>Naïve Bayes</i> , KNN, dan <i>Decision Tree</i> berhasil mencapai akurasi di atas 70% . Hasil terbaik diperoleh oleh <i>Naïve Bayes</i> , yaitu <i>hamming loss</i> 0.1247 .

III. METODE

A. Diagram Sistem



GAMBAR 1
Diagram Alur Sistem

Gambar 1 mengilustrasikan rincian langkah-langkah utama yang dilakukan, mulai dari pengumpulan data hingga evaluasi. Proses dimulai dengan tahap Pengumpulan Data, di mana teks Al-Qur'an yang terdiri dari 6236 ayat diambil dari The Qur'anic Arabic Corpus [15]. Setelah pengumpulan data selesai, tahap selanjutnya adalah Pra-pemrosesan Data, yang mencakup proses tokenisasi, stemming, lemmatization, dan penghapusan kata henti. Langkah-langkah ini bertujuan untuk mempersiapkan data agar lebih terstruktur dan relevan untuk analisis lanjutan. Setelah tahap pra-pemrosesan, dilakukan pembangunan graf kata dengan cara mengidentifikasi hubungan ko-occurrence antara kata-kata dalam setiap ayat, yang direpresentasikan sebagai simpul (node) dan sisi (edges) dalam graf tersebut.

Setelah graf dibangun, langkah berikutnya adalah Ekstraksi Fitur, di mana pengukuran centrality diterapkan untuk mengevaluasi tingkat kepentingan kata-kata dalam graf. Nilai centrality yang diperoleh kemudian digunakan dalam Model Klasifikasi, di mana model Random Forest dan *Naïve Bayes* dilatih untuk mengklasifikasikan ayat-ayat ke dalam delapan kategori tematik. Proses ini diakhiri dengan tahap Evaluasi, di mana kinerja model diukur menggunakan metrik seperti *Hamming loss* untuk menilai efektivitas dari klasifikasi yang dilakukan. Diagram pada Gambar 1. memberikan gambaran yang komprehensif mengenai alur kerja penelitian serta interkoneksi antara setiap langkah dalam rangka mencapai tujuan akhir klasifikasi ayat-ayat Al-Qur'an.

B. Implementasi Sistem

Penelitian ini bertujuan untuk mengidentifikasi serta mengklasifikasikan tema utama yang terdapat dalam Al-Qur'an dengan menggunakan pendekatan pemrosesan bahasa alami (NLP) dan teknik pembelajaran mesin. Dengan menganalisis 6.236 ayat dalam Al-Qur'an, studi ini memadukan metode modern guna memahami pola bahasa dan interaksi antara konsep-konsep yang ada. Proses penelitian meliputi pengumpulan dan pengolahan data, pembuatan graf kata, ekstraksi fitur, pengembangan model klasifikasi, hingga tahap evaluasi, semuanya bertujuan untuk memberikan pemahaman baru mengenai struktur tematik di dalam Al-Qur'an.

1. Pengumpulan Dataset

Peneliti melakukan pengumpulan teks Al-Qur'an yang terdiri dari 6. 236 ayat yang terorganisasi dalam 114 Surah dan 30 Juz. Data yang digunakan adalah teks non-diatik dalam bahasa Arab yang diambil dari *The Qur'anic Arabic Corpus* [15], sebuah proyek internasional yang menyediakan sumber daya linguistik terkait Al-Qur'an. Pemilihan dataset ini didasarkan pada kelengkapan fiturnya, yang memudahkan pemahaman serta diskusi mengenai tata bahasa Arab, sintaksis, dan morfologi setiap kata dalam Al-Qur'an. Meskipun terdapat variasi dalam interpretasi dan pengelompokan ayat oleh para ulama, penelitian ini akan

berfokus pada delapan topik yang telah ditentukan untuk proses pelabelan, yaitu: (1) Akhlak, (2) Alam Akhirat, (3) Alam Dunia, (4) Alam Ghaib, (5) Aqidah, (6) Ilmu (7) Kisah, dan (8) Syariah.

2. Pra-Pemrosesan data

Setelah tahap pengumpulan data selesai, langkah berikutnya adalah pengolahan data, yang mencakup beberapa langkah krusial dalam mempersiapkan data sebelum melakukan analisis lebih lanjut. Pertama-tama, proses tokenisasi dilaksanakan untuk membagi teks menjadi unit-unit lebih kecil yang dikenal sebagai token, yaitu kata-kata. Langkah ini sangat penting untuk memfasilitasi analisis terhadap kata serta hubungan antar kata. Selanjutnya, proses *stemming* dan *lemmatization* diterapkan untuk menghapus prefiks dan sufiks dari kata-kata, sehingga kata-kata dapat direduksi ke bentuk dasarnya. Proses ini berkontribusi dalam mengurangi variasi kata yang tidak perlu dan meningkatkan konsistensi dalam analisis. Selain itu, penghapusan kata henti dilakukan untuk mengeluarkan kata-kata yang sering muncul namun dianggap tidak memiliki makna yang signifikan, seperti konjungsi dan kata ganti. Proses penghilangan kata henti ini dilakukan dengan menggunakan pustaka *NLTK*, yang merupakan alat yang umum digunakan dalam pemrosesan bahasa alami. Gambar 2 Merupakan *word cloud* yang menunjukkan 50 kata dengan tingkat kemunculan tertinggi setelah dilakukan proses penghapusan imbuhan.



GAMBAR 2
Word Cloud untuk 50 kata dengan frekuensi tertinggi

3. Pembuatan Graf

Setelah pengolahan data selesai, langkah selanjutnya adalah membangun graf kata atau *Graph of Words* (GoW), di mana dua jenis graf akan dibuat, satu dengan penghilangan kata henti dan yang lainnya tanpa penghilangan tersebut. Graf ini berfungsi sebagai representasi dari hubungan antar kata, di mana simpul (*nodes*) mewakili kata-kata unik yang ditemukan dalam Al-Qur'an, dan tepi (*edges*) menggambarkan hubungan *ko-occurrence* antara kata-kata tersebut. Dalam penelitian ini, graf yang dibangun bersifat terarah dan berbobot, yang memungkinkan dilakukan analisis lebih mendalam terhadap hubungan kata berdasarkan urutan dan frekuensi kemunculannya. Pemilihan graf terarah memberikan kesempatan kepada peneliti untuk menangkap aliran makna dalam teks, sedangkan sifat berbobot pada graf memberikan bobot yang lebih tinggi pada hubungan yang lebih sering muncul, sehingga mencerminkan pentingnya kata-kata dalam konteks ayat. Setelah grafik kata terbentuk, langkah selanjutnya adalah melakukan pengukuran sentralitas untuk menentukan bobot kepentingan suatu simpul dalam grafik. Semakin tinggi nilai sentralitas, semakin signifikan simpul tersebut dalam konteks GoW.

Berbagai jenis sentralitas ada, namun dalam penelitian ini dibatasi pada tiga jenis: *betweenness*, *closeness*, dan *eigenvector*.

a. *Betweenness Centrality*

Betweenness centrality mengukur seberapa sering suatu *node* berfungsi sebagai penghubung di jalur terpendek antara dua *node* lainnya. *Node* yang memiliki nilai tinggi berperan sebagai penghubung krusial antara bagian-bagian dalam graf [16]. *Betweenness centrality* dapat didefinisikan melalui persamaan berikut:

$$C_B(V_i) = \sum_{V_s \neq V_i \neq V_t \in V, s < t} \frac{\sigma_{st}(V_i)}{\sigma_{st}} \quad (1)$$

Dengan :

$$\sigma_{st} = \text{Jumlah jalur terpendek antara s dan t}$$

$\sigma_{st}(V_i)$ = Jumlah lintasan terpendek antara s dan t yang melewati V_i

b. Betweenness Centrality

Closeness centrality mengukur seberapa dekat suatu node dengan seluruh node lainnya dalam sebuah graf, dengan cara menjumlahkan jarak terpendek dari *node* tersebut ke *node-node* lainnya [17]. Formula untuk menghitung *closeness centrality* adalah:

$$C_c(i) = \frac{n-1}{\sum_{i=1}^n d(i,j)} \quad (2)$$

Dengan:

$$d(i, j) = \text{Jarak antara } i \text{ dengan } j$$

n = Jumlah node

c. Eigenvector Centrality

Eigenvector centrality menilai signifikansi suatu *node* berdasarkan pentingnya *node-node* yang terhubung dengannya. Pendekatan ini memungkinkan analisis yang lebih komprehensif terhadap struktur jaringan dengan memperhatikan tidak hanya kuantitas koneksi, tetapi juga kualitas dari koneksi tersebut [16]. Rumus untuk menghitung *eigenvector centrality* dapat dirumuskan sebagai berikut:

$$E_c(V_i) = \frac{1}{\lambda} \sum_j g_{ij} E_c(V_j) \quad (3)$$

Dengan:

λ = Nilai eigen terbesar dari matriks adjacency graf

g_{ij} = Elemen matriks adjacency graf (V_j) yang terhubung dengan (V_i)

$EE_c(V_j)$ = Eigenvector centrality dari simpul (V_j) yang terhubung dengan (V_i)

4. Ekstraksi Fitur

Proses ekstraksi fitur dalam klasifikasi teks melibatkan konversi data teks menjadi vektor fitur numerik, yang selanjutnya dapat dipakai sebagai input untuk algoritma pembelajaran mesin [18]. Dalam studi ini, setelah penyusunan graf selesai, langkah berikutnya adalah mengekstrak fitur dengan menerapkan metode pengukuran sentralitas yang akan memberikan bobot pada kata-kata dalam graf tersebut. Tiga jenis pengukuran sentralitas yang diterapkan adalah *betweenness*, *closeness*, dan *eigenvector*, yang masing-masing memberikan perspektif unik mengenai pentingnya sebuah kata dalam konteks graf. Selain itu, model ruang vektor digunakan untuk menilai data dengan menggunakan *Term Weight-Inverse Document Frequency* (TW-IDF) untuk eksperimen utama. Nilai sentralitas dan

bobot yang dihasilkan selanjutnya digunakan sebagai fitur dalam model klasifikasi multilabel, memungkinkan peneliti untuk mengidentifikasi kata-kata yang paling relevan bagi setiap kategori tematik yang telah ditetapkan.

a. Term Frequency-Inverse Document Frequency

TF-IDF (*Term Frequency-Inverse Document Frequency*) merupakan suatu metode penghitungan bobot yang diterapkan dalam analisis teks guna menilai pentingnya suatu kata dalam konteks dokumen tertentu dibandingkan dengan keseluruhan kumpulan dokumen (korpus). Metode ini menghitung nilai dengan mengalikan dua elemen utama, *Term Frequency* (TF), yang menunjukkan frekuensi kemunculan sebuah kata dalam dokumen tertentu, dan *Inverse Document Frequency* (IDF), yang mencerminkan seberapa jarang kata tersebut muncul di seluruh dokumen dalam korpus. Kata-kata yang muncul dengan frekuensi tinggi dalam satu dokumen tetapi jarang ditemukan di dokumen lain akan mendapatkan bobot yang tinggi, karena dianggap relevan untuk dokumen tersebut [19]. TF-IDF berperan penting dalam menekankan kata-kata yang lebih signifikan dibandingkan dengan kata-kata umum seperti "dan" atau "adalah."

b. Term Weight-Inverse Document Frequency

TW-IDF (*Term Weight-Inverse Document Frequency*) merupakan suatu inovasi dari TF-IDF yang menambahkan elemen bobot pada kata, yakni *Term Weight* (TW), untuk memberikan representasi yang lebih relevan dengan konteks. Bobot ini dapat diperoleh dari berbagai atribut tambahan, seperti ukuran sentralitas dalam graf (misalnya, *eigenvector* atau *betweenness centrality*) [20]. Dengan kata lain, TW-IDF tidak hanya memperhatikan frekuensi kata dalam dokumen dan keunikan kata dalam korpus, tetapi juga mengintegrasikan informasi tambahan yang berkaitan dengan signifikansi kata dalam jaringan atau struktur tertentu. Pendekatan ini menawarkan bobot yang lebih fleksibel dalam analisis teks yang berbasis graf.

5. Perancangan Model

Nilai sentralitas yang diekstraksi dari graf digunakan sebagai fitur dalam proses klasifikasi multilabel. Dalam penelitian ini, terdapat dua model utama yang diterapkan, yaitu *Random Forest* dan *Naive Bayes*. Model *Random Forest* digunakan untuk mengungkap hubungan non-linear dalam data melalui pendekatan berbasis pohon keputusan, yang telah terbukti efektif dalam beragam aplikasi klasifikasi. Sementara itu, *Naive Bayes* dipilih karena kesederhanaan dan efisiensinya dalam mengelola data teks berbasis probabilitas, penelitian [21] menunjukkan bahwa model ini dapat memberikan tingkat akurasi yang tinggi dalam klasifikasi multilabel. Sebelum melakukan pelatihan, dataset dibagi menjadi data pelatihan dan data pengujian sehingga model dapat belajar dari data yang ada dan diuji menggunakan data baru.

6. Klasifikasi Multilabel

Sebelum melakukan klasifikasi, dataset dibagi menjadi dua bagian, set pelatihan dan set pengujian dengan proporsi 70% untuk pelatihan dan 30% untuk pengujian. Pembagian ini bertujuan untuk melatih model agar mampu

memprediksi label dengan akurasi tinggi. Mengingat adanya distribusi label yang tidak seimbang dalam dataset, diterapkan metode *oversampling* seperti SMOTE (*Synthetic Minority Over-sampling Technique*) untuk meningkatkan representasi kelas minoritas tanpa mengorbankan data yang berharga. Dengan langkah ini, model dapat dilatih dengan lebih efektif untuk mengenali dan mengklasifikasikan semua label yang tersedia. Setelah proses pelatihan, model yang telah dilatih dievaluasi menggunakan *Hamming loss*, yang berfungsi untuk mengukur tingkat kesalahan prediksi dalam klasifikasi multilabel. Penelitian ini mengaplikasikan dua model utama, yaitu *Random Forest* dan *Naive Bayes*.

7. Evaluasi

Proses evaluasi dalam klasifikasi multilabel ini menggunakan metrik *Hamming Loss*, yang mengukur proporsi kesalahan prediksi label terhadap total label yang tersedia. Semakin rendah nilai *Hamming Loss*, semakin baik kinerja model dalam mengklasifikasikan label yang benar.

Evaluasi dilakukan dengan membandingkan berbagai metode ekstraksi fitur berdasarkan *centrality*, seperti *closeness*, *betweenness*, dan *eigenvector*, yang dipadukan dengan algoritma klasifikasi seperti *Random Forest* dan *Naive Bayes*. Selain itu, penggunaan *Synthetic Minority Over-sampling Technique* (SMOTE) juga diuji untuk menangani ketidakseimbangan data pada kelas minoritas, dengan hasil yang menunjukkan dampak yang beragam terhadap akurasi. Evaluasi ini sangat penting untuk menentukan efektivitas sistem dalam mengkategorikan ayat-ayat Al-Qur'an ke dalam delapan tema tematik, serta untuk mengidentifikasi metode terbaik guna meningkatkan akurasi klasifikasi.

IV. HASIL DAN PEMBAHASAN

A. Hasil Pengujian

Hasil pengujian utama menunjukkan bahwa algoritma *Random Forest* menghasilkan performa terbaik dengan nilai *Hamming Loss* terendah. Dalam analisis tersebut, nilai *Hamming Loss* untuk sentralitas kedekatan (*closeness centrality*) tanpa penghilangan kata henti untuk *Random Forest* mencapai **0.1631**, dan *Naive Bayes* berada pada 0.1888. Setelah penerapan penghilangan kata henti, nilai *Hamming Loss* untuk *Random Forest* menjadi 0.1655 dan *Naive Bayes* menjadi 0.1891. Temuan ini menunjukkan bahwa penghilangan kata henti dapat sedikit mengurangi performa model.

Selanjutnya, pada centralitas perantara (*betweenness centrality*), analisis menunjukkan bahwa tanpa penghilangan kata henti, *Random Forest* memiliki *Hamming Loss* sebesar 0.1639, dan *Naive Bayes* mencatat **0.1857**. Setelah penerapan penghilangan kata henti, nilai *Hamming Loss* untuk *Random Forest* dan *Naive Bayes* menunjukkan perubahan yang minimal, yaitu **0.1631** dan 0.1873, berturut-turut. Hal ini mengindikasikan bahwa penghilangan kata henti sedikit meningkatkan kinerja model, meskipun tetap tidak dapat mengungguli sentralitas kedekatan tanpa penghilangan kata henti.

Pada centralitas *eigenvector* (*eigenvector centrality*), tanpa penghilangan kata henti, *Random Forest* mencatat nilai *Hamming Loss* sebesar **0.1631**, dan *Naive Bayes* memperoleh 0.1885. Setelah penerapan penghilangan kata henti, nilai *Hamming Loss* untuk *Random Forest* menjadi 0.1633 dan

Naive Bayes menjadi 0.1892. Temuan ini menunjukkan bahwa penghilangan kata henti tidak memberikan peningkatan yang signifikan pada sentralitas *eigenvector*.

Sebagai perbandingan, penelitian [11] mencatat nilai *Hamming Loss* sebesar **0.1615** untuk CNN dan **0.1694** untuk RNN (*Recurrent Neural Network*) tanpa penghilangan kata henti. Kedua nilai ini menunjukkan performa yang lebih rendah dibandingkan dengan hasil penelitian ini yang menggunakan sentralitas *eigenvector*.

Secara keseluruhan, hasil penelitian ini menunjukkan bahwa penggunaan *closeness centrality* tanpa menghilangkan kata henti menghasilkan performa terbaik pada model *Random Forest*. Hal ini sejalan dengan temuan dalam penelitian [11], yang juga mengindikasikan keunggulan serupa pada model CNN dan RNN dengan penggunaan sentralitas yang sama. Untuk lebih detailnya dapat dilihat pada Tabel 2.

TABEL 2
HASIL PENGUJIAN MODEL KLASIFIKASI DENGAN BERBAGAI JENIS CENTRALITY (HAMMING LOSS)

Centrality	Model	Tanpa Stopwords (Pada penelitian ini)	Menggunakan Stopwords (Pada Penelitian ini)	Tanpa Stopwords (Pada penelitian [11])	Menggunakan Stopwords (Pada penelitian [11])
Betweenness	CNN	-	-	0.2175	0.2208
	RNN	-	-	0.2162	0.2073
	Random Forest	0.1639	0.1631	-	-
	Naive Bayes	0.1857	0.1873	-	-
Closeness	CNN	-	-	0.1615	0.1687
	RNN	-	-	0.1694	0.1767
	Random Forest	0.1631	0.1655	-	-
	Naive Bayes	0.1888	0.1891	-	-
Eigenvector	CNN	-	-	0.1828	0.1845
	RNN	-	-	0.1842	0.1848
	Random Forest	0.1631	0.1633	-	-
	Naive Bayes	0.1885	0.1892	-	-

Keterangan:

- Font Bold Warna Hitam : Tingkat *hamming loss* terendah menggunakan Algoritma *Naive Bayes* pada penelitian ini.
- Font Bold Warna Biru : Tingkat *hamming loss* terendah menggunakan Algoritma *Random Forest* pada penelitian ini.
- Font Bold Warna Hijau : Tingkat *hamming loss* menggunakan Algoritma CNN dan RNN pada penelitian lain.

B. Analisis Hasil Pengujian

Berdasarkan hasil pengujian yang dilaksanakan, *betweenness centrality* menunjukkan performa terbaik dibandingkan dengan berbagai jenis centrality yang diuji, di mana model *Random Forest* dengan penggunaan kata henti sebesar **0.1631** dan *Naive Bayes* tanpa penggunaan kata henti sebesar **0.1857**. Temuan ini mengindikasikan bahwa *betweenness centrality* lebih efektif dalam mengidentifikasi hubungan antar kata dalam teks Al-Qur'an tanpa menghilangkan kata-kata henti yang mungkin memiliki relevansi penting dalam konteks semantik. Penggunaan penghilangan kata henti justru berujung pada peningkatan *Hamming Loss* dalam beberapa kasus, khususnya pada *closeness centrality*, yang menunjukkan bahwa kata-kata henti memiliki kontribusi signifikan dalam menjaga hubungan antar kata yang relevan dalam teks.

Secara keseluruhan, penggunaan *betweenness centrality* tanpa penghilangan kata henti, yang dikombinasikan dengan

model *Random Forest*, merupakan kombinasi paling efektif untuk mengklasifikasikan ayat-ayat Al-Qur'an ke dalam kategori tematik pada penelitian ini, memberikan pemahaman yang lebih dalam tentang struktur semantik teks, dan menjadi dasar untuk pengembangan sistem klasifikasi teks dalam kajian literatur Islam di masa depan.

V. KESIMPULAN

Berdasarkan hasil pengujian dan analisis yang telah dilaksanakan, beberapa poin penting dapat disimpulkan. *Betweenness centrality* dengan penerapan penghilangan kata henti menunjukkan performa terbaik dalam klasifikasi ayat-ayat Al-Qur'an, dengan dua algoritma yang memiliki nilai *Hamming Loss* terendah di antara semua jenis centrality yang diuji. Model *Random Forest* memperoleh hasil terbaik dengan nilai *Hamming Loss* sebesar 0.1631 pada *closeness centrality* tanpa penghilangan kata henti, *betweenness centrality* dengan penghilangan kata henti, dan *eigenvector centrality* tanpa penghilangan kata henti. Namun, penghilangan kata henti tidak selalu meningkatkan performa; dalam beberapa kasus, hal ini justru menyebabkan peningkatan nilai *Hamming Loss*, terutama pada *closeness centrality*.

Di masa mendatang, diperlukan eksplorasi lebih lanjut mengenai peran kata henti dalam teks Al-Qur'an dengan mempertimbangkan kedekatannya terhadap kata-kata signifikan. Pengujian menggunakan dataset lain, seperti teks dari literatur Islam atau dalam bahasa yang berbeda, juga dapat memberikan wawasan lebih mendalam. Selain itu, penerapan teknik graph mining lainnya, seperti *graph neural networks* (GNN), patut dipertimbangkan untuk meningkatkan representasi graf.

REFERENSI

- [1] A. Rofiqul Muslikh, I. Akbar, D. Rosal Ignatius Moses Setiadi, and H. Md Mehedul Islam, "Multi-label Classification of Indonesian Al-Quran Translation based CNN, BiLSTM, and FastText." [Online]. Available: <https://quran.kemenag.go.id>.
- [2] F. Akbar Nugroho Ilmu Komputer, "PENGUNAAN TEXT MINING DALAM ANALISIS BIG DATA: TREN TERBARU," 2024.
- [3] A. Saputra et al., "Big Data," 2022.
- [4] F. Hakim, A. Fadlillah, M. Nafiur Rofiq, and U. Al Falah Assunniah Kencong Jember, "Artificial Intellegence (AI) dan Dampaknya Dalam Distorsi Pendidikan Islam," vol. 13, no. 1, 2024, doi: 10.54437/juw.
- [5] K. Gufran Mursyid and M. Awaliyah, "MAKKIYAH DAN MADANIYAH DALAM AL-QUR'AN," 2021.
- [6] "L. N. Hakim, Visualisasi Tematik Al-Qur'an Berbasis Knowledge Graph, Skripsi, Department of Informatics Engineering, Universitas Islam Riau, Pekanbaru, Indonesia, 2019."
- [7] B. Yusuf et al., "Analisa Topik Pendidikan Dalam Al-Quran dengan Pendekatan Text Mining," *Serambi Engineering*, vol. VI, no. 1, 2021, [Online]. Available: <https://quran.kemenag.go.id/>.
- [8] S. Delimasari and K. Kusri, "Komparasi Algoritma Machine Learning Untuk Menganalisis Sentimen Ulasan Pada Aplikasi Digital Korlantas Polri," *G-Tech: Jurnal*

Teknologi Terapan, vol. 8, no. 4, pp. 2411–2419, Oct. 2024, doi: 10.70609/gtech.v8i4.5089.

[9] R. Diki Nugraha, “Call for papers dan Seminar Nasional Sains dan Teknologi Ke-2 2023 Fakultas Teknik, Universitas Pelita Bangsa,” vol. 2, no. 1, 2023.

[10] E. Supriyati and M. Iqbal, “PENGUKURAN SIMILARITY TEMA PADA JUZ 30 AL QUR’AN MENGGUNAKAN TEKS KLASIFIKASI,” *Jurnal SIMETRIS*, vol. 9, no. 1, 2018, [Online]. Available: <https://translate.google.co.id/>

[11] F. R. Muflihah, K. M. Lhaksana, and M. A. Bijaksana, “Centrality-Based Multilabel Neural Networks Classification of Qur’an Verse Topics,” *2024 International Conference on Data Science and Its Applications (ICoDSA)*, pp. 469–474, 2024, doi: 10.1109/ICODSA.2024

[12] F. Yulianto, K. M. Lhaksana, and D. T. Murdiansyah, “Classifying Quranic Verse Topics using Word Centrality Measure,” *Jurnal RESTI (Rekayasa Sistem dan Teknologi Informasi)*, vol. 5, no. 3, pp. 594–601, Jun. 2021, doi: 10.29207/resti.v5i3.3171.

[13] M. I. Rahman, N. A. Samsudin, A. Mustapha, and A. Abdullahi, “Comparative analysis for topic classification in Juz Al-Baqarah,” *Indonesian Journal of Electrical Engineering and Computer Science*, vol. 12, no. 1, pp. 406–411, Oct. 2018, doi: 10.11591/ijeecs.v12.i1.pp406-411.

[14] Achmad Salim Aiman, Kemas Muslim Lhaksana, and Jondri, “Topic Classification of Quranic Verses in English Translation Using Word Centrality Measurement,” *Jurnal RESTI (Rekayasa Sistem dan Teknologi Informasi)*, vol. 6, no. 5, pp. 803–809, Oct. 2022, doi: 10.29207/resti.v6i5.4358.

[15] “Quranic Arabic Corpus - Data Download.” Accessed: Dec. 06, 2024. [Online]. Available: <https://corpus.quran.com/download/>

[16] Y. Dwi and P. Ariyanti, “Analisis Centrality Aktor pada Penyebaran Informasi Kuliner di Media Sosial dengan menggunakan Social Network Analysis.” [Online]. Available: <http://e-journal.ivet.ac.id/index.php/jsitee>

[17] H. Tuhuteru and A. Iriani, “Analisis Kolaborasi Penelitian Ilmiah Dosen Fakultas X dengan Social Network Analysis (SNA),” *Jurnal Teknik Informatika dan Sistem Informasi*, vol. 4, no. 1, Apr. 2018, doi: 10.28932/jutisi.v4i1.758.

[18] L. Efrizoni, S. Defit, M. Tajuddin, and A. Anggrawan, “Komparasi Ekstraksi Fitur dalam Klasifikasi Teks Multilabel Menggunakan Algoritma Machine Learning,” *MATRIK: Jurnal Manajemen, Teknik Informatika dan Rekayasa Komputer*, vol. 21, no. 3, pp. 653–666, Jul. 2022, doi: 10.30812/matrik.v21i3.1851.

[19] M. Nurjannah, I. Fitri Astuti, and D. Program Studi, “PENERAPAN ALGORITMA TERM FREQUENCY-INVERSE DOCUMENT FREQUENCY (TF-IDF) UNTUK TEXT MINING Mahasiswa S1 Program Studi Ilmu Komputer FMIPA Universitas Mulawarman 2,3),” 2013.

[20] P. Diseminasi and F. Genap, “Akurasi dan Presisi Pengklasifikasian Abstrak Paper Informatika Menggunakan TF-IDF dan Multiclass Support Vector Machine (SVM).”

[21] K. I. Gunawan and J. Santoso, “Multilabel Text Classification Menggunakan SVM dan Doc2Vec Classification Pada Dokumen Berita Bahasa Indonesia,” *Journal of Information System, Graphics, Hospitality and Technology*, vol. 3, no. 01, pp. 29–38, Apr. 2021, doi: 10.37823/insight.v3i01.126.