

Pengelompokan Pernyataan Misogini pada Media Sosial Berbahasa Indonesia Menggunakan TF-IDF dan Algoritma *K-Means*

Umaya Ramadhani Putri Nasution ^{1*}, Ikhwanuddin Nasution ², Lia Silviana ³, Fanindia Purnamasari ³, Marischa Elveny ³, Anandhini Medianty Nababan ³, Stephani Uli Basa Silitonga ³

¹ Fakultas Ilmu Komputer dan Teknologi Informasi, Teknologi Informasi, Universitas Sumatera Utara, Medan, Indonesia

² Fakultas Ilmu Budaya, S-2 Penciptaan dan Pengkajian Seni, Universitas Sumatera Utara, Medan, Indonesia

³ Fakultas Ilmu Komputer dan Teknologi Informasi, Teknologi Informasi, Universitas Sumatera Utara, Medan, Indonesia

INFORMASI ARTIKEL

Diterima Redaksi: 02 Juni 2026

Revisi Akhir: 10 Juli 2026

Diterbitkan Online: 22 Juli 2026

KATA KUNCI

Analisis Teks
Clustering
K-Means
Media Sosial
Misogini

KORESPONDENSI (*)

Phone: +62 813-1704-0340

E-mail: umaya.nst@usu.ac.id

A B S T R A K

Penelitian ini bertujuan untuk mengidentifikasi pola tersembunyi pada komentar media sosial berbahasa Indonesia yang berkaitan dengan misogini menggunakan pendekatan clustering. Dataset yang digunakan terdiri atas komentar yang dikumpulkan dari beberapa platform media sosial. Tahapan penelitian meliputi preprocessing teks, ekstraksi fitur menggunakan metode Term Frequency-Inverse Document Frequency (TF-IDF), serta pengelompokan data menggunakan algoritma K-Means. Penentuan jumlah kluster optimal dilakukan melalui evaluasi menggunakan Silhouette Score, Davies-Bouldin Index, dan Calinski-Harabasz Index, yang menghasilkan jumlah kluster optimal sebanyak lima kluster. Hasil analisis top terms menunjukkan bahwa setiap kluster memiliki karakteristik semantik yang berbeda, meliputi objektifikasi seksual eksplisit, diskursus peran gender, benevolent sexism, interaksi relasional, dan objektifikasi fisik yang dinormalisasi. Hasil penelitian menunjukkan bahwa komentar yang berkaitan dengan misogini tidak bersifat homogen, tetapi membentuk beberapa kelompok berdasarkan pola bahasa dan konteks semantik yang berbeda. Temuan ini menunjukkan bahwa kombinasi representasi fitur berbasis TF-IDF dan algoritma K-Means mampu mengungkap struktur semantik laten pada komentar media sosial yang tidak dapat direpresentasikan secara memadai melalui kategorisasi biner misogini dan non-misogini.

PENDAHULUAN

Hingga saat ini, banyak masyarakat Indonesia yang menganut paham patriarki. Hal ini menyebabkan misogini berkembang secara luas seiring dengan perkembangan ideologi tersebut sampai lahirnya seorang misoginis. Seorang misoginis memiliki sikap yang memandang bahwa perempuan merupakan pihak yang terlibat dalam konflik, dan definisi misoginis sendiri didasarkan pada kebencian terhadap perempuan, dimana banyak orang percaya bahwa perempuan pantas menderita dan dilecehkan, sehingga perempuan dihina dan diremehkan [1].

Berdasarkan laporan oleh *National Democratic Institute* [2], dinyatakan bahwa perempuan di Indonesia menjadi yang paling sering mengalami misogini tersebut. Dikarenakan maraknya penggunaan sosial media dimana-mana dan akun yang paling banyak menggunakannya adalah perempuan, penggunaan sosial media tersebut dijadikan alat untuk dilakukannya misogini. Oleh karena itu, perlu ada tindak lanjut terkait misogini tersebut yang telah merugikan perempuan.

Sudah cukup banyak pendekatan yang dilakukan terkait misogini. Ada beberapa penelitian yang telah dilakukan seperti “Identifikasi pernyataan Misogini berbahasa Indonesia Berdasarkan Komentar Youtube Menggunakan Glove Embedding dan Random Forest Classisfier” [3]. Dimana dalam penelitian ini berhasil mengidentifikasi sebesar 92,5 % dalam komentar youtube yang dijadikan sebagai sampel data dinyatakan banyak yang mengandung misogini. Kemudian pada penelitian lain yang berjudul “*Misogynous Text Classification Using SVM and LSTM*” [4], dimana penelitian ini menggunakan dataset berbahasa India dan Inggris. Membandingkan hasil klasifikasi dari dua bahasa dataset yang berbeda. Penelitian lain dengan judul “*Identification of Misogyny on Social Media in Indonesia Using Bidirectional Encoder Representations From Transformers (BERT)*” [5]. Pada penelitian tersebut diambil dari dua sosial media yaitu instagram dan youtube dan menggunakan IndoBERT sebagai model penelitiannya. Dan masih banyak beberapa penelitian lainnya terkait misogini dalam hal indentifikasi dan klasifikasi.

Pada penelitian-penelitian sebelumnya, komentar media sosial berlabel misogini dan non-misogini umumnya dianalisis menggunakan pendekatan identifikasi dan klasifikasi. Namun, pendekatan tersebut belum mampu mengungkap struktur semantik laten serta variasi bentuk ekspresi misogini yang tersembunyi dalam data. Oleh sebab itu, diperlukan pendekatan unsupervised learning untuk mengeksplorasi pola-pola alami pada komentar media sosial tanpa bergantung pada label yang telah ditentukan sebelumnya. Berdasarkan permasalahan tersebut, penelitian ini melakukan pengelompokan terhadap pernyataan misogini pada media sosial berbahasa Indonesia menggunakan pendekatan clustering untuk mengidentifikasi karakteristik dan pola bahasa yang terbentuk.

TINJAUAN PUSTAKA

Clustering

Clustering adalah sebuah metode analisis data yang mengelompokkan objek-objek berdasarkan kemiripannya [6]. Proses analisis dari klasterisasi disini melihat tingkat kemiripan dari atribut-atribut data yang dimiliki yang sebelumnya tidak diketahui label kelasnya, namun di beberapa kejadian bisa menjadi mungkin terhadap data dengan label kelas yang diketahui. Dalam proses klasterisasi, tingkat kemiripan diukur dengan memasukkan objek yang sama ke dalam klaster yang sama dan objek yang berbeda ke dalam klaster yang memiliki perbedaan yang sama [7]. Evaluasi perbedaan ini dilakukan dengan menggunakan nilai atribut objek dan metrik jarak. Salah satu algoritma clustering adalah K-Means. Pembahagian data saat ini menjadi satu atau lebih cluster atau grup adalah hasil dari metode pengelompokan data non-hierarki yang disebut K-Means [8].

K-Means

K-means adalah algoritma yang terdiri dari dua bagian yang berbeda dan digunakan untuk memecahkan berbagai masalah pengelompokan data untuk dataset besar. Algoritma ini terdiri dari jumlah pusat yang dipilih secara acak kemudian diperbaiki, dan k adalah jumlah pusat yang dipilih secara acak [9]. Algoritma K-Means digunakan dalam data mining untuk mengelompokkan data berdasarkan jarak daripada berdasarkan variabel [10].

Dalam penelitian ini, metode Term Frequency–Inverse Document Frequency (TF-IDF) digunakan untuk menggambarkan setiap dokumen dalam bentuk vektor fitur. Representasi yang dihasilkan kemudian dimasukkan ke dalam algoritma K-Means untuk melakukan pengelompokan dokumen berdasarkan kedekatan antar dokumen.

Algoritma K-Means mengukur tingkat kedekatan antar dokumen. Ini dilakukan dengan menggunakan perhitungan jarak geometris. Untuk mengetahui jarak antara dua dokumen atau data ke-i terhadap pusat klaster ke-j, gunakan persamaan berikut:

Untuk dua vektor $X = (x_1, x_2, \dots, x_n)$ dan $Y = (y_1, y_2, \dots, y_n)$

$$d(X, Y) = \sqrt{\sum_{i=1}^n (X_i - Y_i)^2} \quad (1)$$

Keterangan:

n = jumlah dimensi (fitur teks)

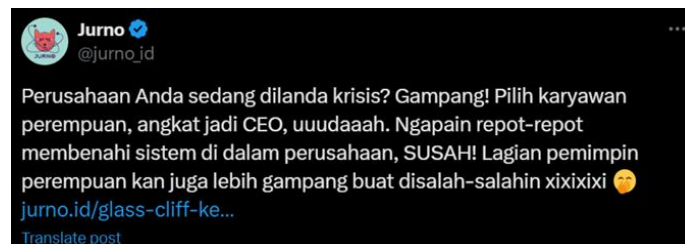
X_i = nilai fitur ke-i dari dokumen pertama

Y_i = nilai fitur ke- i dari dokumen kedua

Pernyataan Misogini

Dalam bahasa Yunani misogini diartikan dari dua kata yaitu “misein” dengan arti benci dan “gyne” dengan arti perempuan [11]. Adapun dalam arti luas misogini adalah tindakan yang melecehkan perempuan dalam berbagai hal dan tindakannya bisa secara langsung maupun tidak langsung atau dalam kehidupan virtual. Misalnya contoh tindakan langsung adalah melecehkan perempuan secara fisik, sedangkan contoh dalam kehidupann virtual misalnya satu pernyataan dalam sosial media yang ada, yaitu “Itu nenennya tepos sekali”.

Berikut contoh pernyataan misogini yang diambil dari sosial media bisa dilihat pada Gambar 1 di bawah.



Gambar 1. Pernyataan misogini

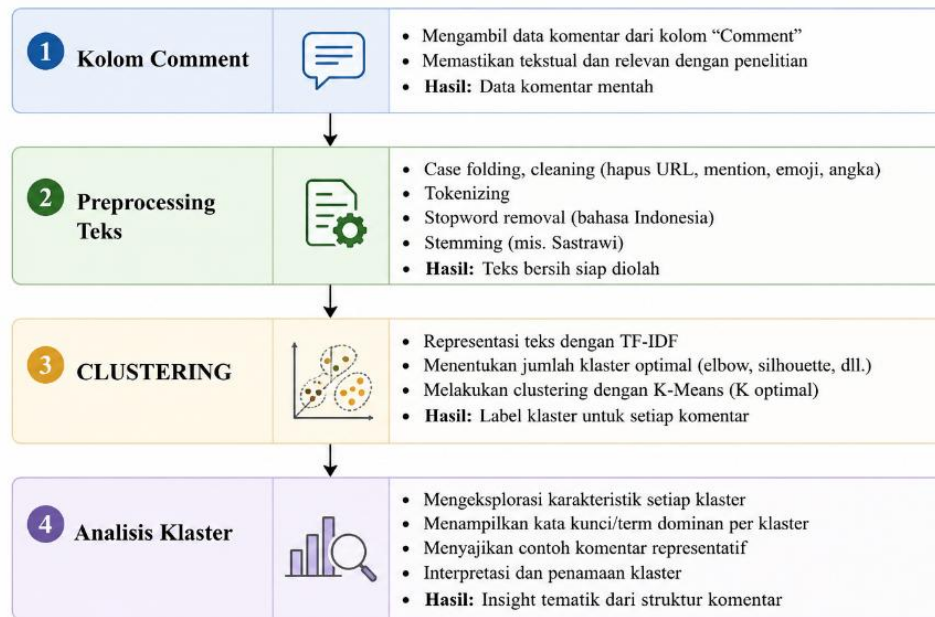
METODOLOGI

Data yang digunakan pada penelitiann ini adalah pernyataan yang diambil dari sosial media berupa komentar, dimana sosial media yang digunakan adalah instagram, youtube dan X. Proses pengambilan komentar melalui situs Instagram dilakukan dengan menggunakan library Python Selenium, yang memiliki kemampuan untuk otomatisasi penggunaan situs sehingga sesuai untuk membuat catatan komentar, sedangkan untuk melakukan proses pengambilan komentar di YouTube. Penelitian ini menggunakan YouTube Data API dimana memberikan akses ke akun channel, daftar putar, dan video, yang memungkinkan untuk mengambil tanggapan. Kemudian, pada X, pengambilan komentar dilakukan menggunakan beberapa kata kunci tertentu saat mengumpulkan data menggunakan library tweet-harvest. Setelah itu, semua pernyataan yang dipilah dari sosial media tersebut dikumpulkan ke dalam file.cs. Total 4.000 pernyataan yang akan digunakan dalam penelitian ini.

Tabel 1. Contoh pernyataan misogini

No	Pernyataan
1	RT @tirz_fabz: Jadi intinya laki2 memang lbh dipercaya utk jd pemimpin dibandingkan perempuan
2	Kalau saja dulu ibumu terpapar virus pemikiran feminis Child Free niscaya kamu sekarang tak akan ada Gosah sok an Child Free padahal aslinya males ngurus anak Awowkwowkwo #feminisindonesia #childfree #ngakakkocak
3	Pinter bacot ngide doang nih cewe kerjaannya jajan mulu keluyuran kek lonte
4	@AiraNtieReal Artis yg sedikit job nya kalo ngak jdi lonte pejabat ngak bakal 23ungga apartemen dan mobil.. Pejabat kalo ngak ngelonte ngak afdol
5	@RadioElshinta Supporter bola club Indonesia yg begini emg lahir dari 23ungg lonte yg memeknya dipake ribuan kontol Pas lahir dulu daripada dibesarin mending kasih ke lele

Gambar 2 menunjukkan tahapan penelitian yang terdiri dari proses pengumpulan data komentar media sosial, preprocessing teks, ekstraksi fitur menggunakan TF-IDF, proses clustering menggunakan algoritma K-Means, evaluasi jumlah kluster optimal, serta reduksi dimensi menggunakan PCA untuk visualisasi hasil pengelompokan.



Gambar 2. Alur metodologi clustering pernyataan misogini

Setelah proses preprocessing dilakukan, data teks yang telah dibersihkan masih berupa data tidak terstruktur dalam bentuk kumpulan kata. Agar dapat diproses menggunakan algoritma clustering, teks perlu diubah menjadi representasi numerik. Pada penelitian ini, ekstraksi fitur dilakukan menggunakan metode Term Frequency–Inverse Document Frequency (TF-IDF).

TF-IDF adalah metode pembobotan kata yang digunakan untuk mengukur tingkat kepentingan suatu kata dalam sebuah dokumen berdasarkan frekuensi kemunculannya pada dokumen dan tingkat kemunculannya pada seluruh koleksi dokumen. Metode ini memberikan bobot yang lebih tinggi terhadap kata-kata yang memiliki karakteristik tertentu, sementara pengaruh kata-kata yang sering muncul di banyak dokumen dikurangi.

Nilai Term Frequency (TF) dan Inverse Document Frequency (IDF) digabungkan untuk menghasilkan nilai TF-IDF. Nilai TF menunjukkan frekuensi kemunculan suatu kata dalam dokumen, sedangkan nilai IDF menunjukkan tingkat keunikan kata dalam dokumen secara keseluruhan.

$$\text{TFIDF}(t,d) = \text{TF}(t,d) \times \text{IDF}(t) \quad (2)$$

Keterangan:

$\text{TFIDF}(t,d)$ = bobot kata t pada dokumen d
 $\text{TF}(t,d)$ = frekuensi kemunculan kata t dalam dokumen d
 $\text{IDF}(t)$ = nilai inverse document frequency

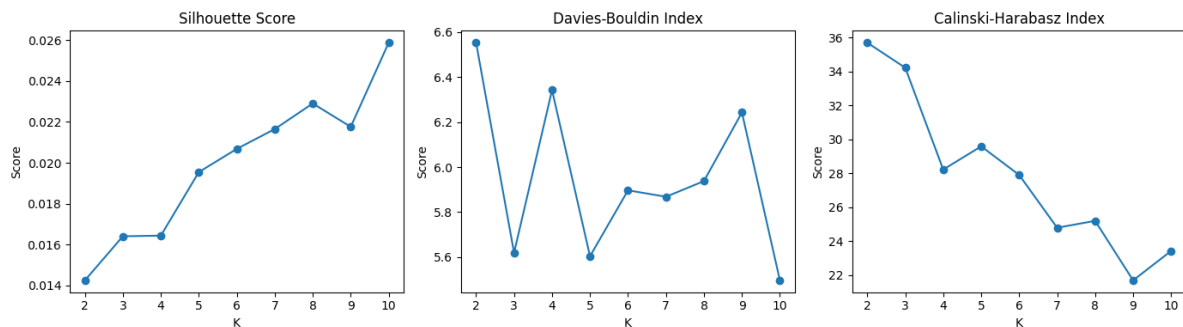
Proses TF-IDF menghasilkan matriks numerik yang memiliki dimensi dokumen $\times$ fitur kata. Matriks ini kemudian digunakan sebagai input algoritma K-Means untuk melakukan proses pengelompokan yang didasarkan pada kemiripan karakteristik semantik antar dokumen.

Hasil kluster akan mengungkap pola laten dimana pernyataan misogini bisa terdiri dari beberapa sub kelompok, bisa menghasilkan pola bahasa misogini implisit vs eksplisit, dan kluster non-misogini yang mirip misogini (*borderline*) bisa ditemukan.

HASIL DAN PEMBAHASAN

Dalam penentuan jumlah kluster optimal digunakan 3 (tiga) metode yaitu, Calinski–Harabasz Index, Silhouette Score dan Davies–Bouldin Index. Ketiga metrik ini memiliki fitur evaluasi yang berbeda. Calinski–Harabasz Index menilai rasio dispersi antar kluster terhadap dispersi dalam kluster, Silhouette Score menilai tingkat kohesi dan separasi antar kluster,

dan Davies-Bouldin Index menilai rata-rata kemiripan antar kluster (nilai yang lebih rendah menunjukkan kualitas yang lebih baik). Berikut adalah hasil metrik dari ketiga metode tersebut.



Gambar 3. Hasil metrik penentuan k terbaik

Berdasarkan hasil evaluasi, nilai Silhouette Score tertinggi diperoleh pada konfigurasi $k=10$. Namun, konfigurasi tersebut menghasilkan jumlah kluster yang relatif banyak sehingga interpretasi karakteristik setiap kelompok menjadi lebih kompleks. Pada data teks media sosial, peningkatan jumlah kluster tidak selalu menunjukkan struktur semantik yang lebih baik, karena dapat menyebabkan fragmentasi terhadap kelompok komentar yang memiliki pola makna serupa.

Sementara itu, nilai Calinski-Harabasz Index tertinggi diperoleh pada $k=2$. Meskipun secara matematis konfigurasi tersebut menunjukkan pemisahan antar kelompok yang baik, jumlah dua kluster cenderung menghasilkan pembagian yang terlalu umum. Konfigurasi tersebut berpotensi hanya memisahkan data menjadi kelompok besar, seperti komentar dengan kecenderungan misogini dan komentar non-misogini, sehingga kurang sesuai dengan tujuan penelitian yang berfokus pada eksplorasi variasi pola bahasa misogini yang lebih spesifik.

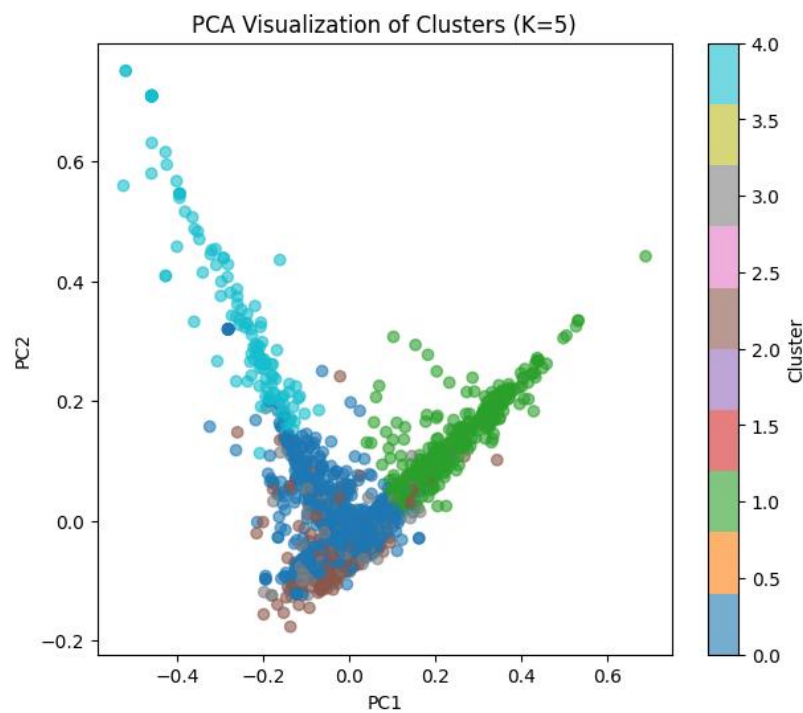
Pada konfigurasi $k=5$, hasil evaluasi menunjukkan keseimbangan antara kualitas pemisahan kluster dan kemudahan interpretasi. Nilai Davies-Bouldin Index menunjukkan pola penurunan dan mencapai titik yang relatif rendah pada konfigurasi ini, yang menunjukkan bahwa jarak antar kluster dan tingkat homogenitas dalam kluster berada pada kondisi yang lebih baik dibandingkan beberapa konfigurasi lainnya. Selain itu, jumlah lima kluster memungkinkan analisis karakteristik semantik setiap kelompok secara lebih mendalam tanpa menyebabkan fragmentasi berlebihan.

Dengan mempertimbangkan ketiga metrik evaluasi serta tujuan penelitian untuk menemukan variasi pola ekspresi misogini pada media sosial, konfigurasi $k=5$ dipilih sebagai jumlah kluster optimal. Pemilihan ini tidak hanya didasarkan pada nilai tertinggi dari satu metrik tertentu, tetapi melalui pertimbangan keseimbangan antara kualitas clustering, interpretabilitas hasil, dan kesesuaian dengan tujuan analisis.

Karena setiap kata unik dalam dataset akan menjadi satu fitur, matriks fitur hasil ekstraksi TF-IDF biasanya memiliki dimensi yang sangat besar. Proses interpretasi visual hasil clustering dapat menjadi lebih sulit karena kondisi ini. Oleh karena itu, untuk menunjukkan hasil pengelompokan, penelitian ini menggunakan Principal Component Analysis (PCA) sebagai metode reduksi dimensi.

Setelah proses clustering, PCA digunakan untuk memetakan hasil pengelompokan ke dalam ruang dua dimensi (2D), sehingga pola distribusi setiap kluster dapat dilihat secara visual pada Gambar 4. PCA adalah teknik reduksi dimensi yang mempertahankan informasi atau variasi terbesar dari data saat mengubah data yang lebih besar ke data yang lebih kecil.

Penggunaan PCA pada penelitian ini hanya bertujuan untuk kebutuhan visualisasi dan interpretasi hasil clustering, bukan sebagai proses pengurangan fitur sebelum algoritma K-Means dilakukan.



Gambar 4. Hasil kluster dari PCA

Visualisasi hasil clustering menggunakan PCA pada Gambar 4 menunjukkan distribusi lima kluster dalam ruang dua dimensi berdasarkan representasi fitur TF-IDF. Hasil visualisasi memperlihatkan bahwa beberapa kluster masih mengalami tumpang tindih (overlapping), terutama pada kelompok yang memiliki kemiripan karakteristik semantik. Kondisi ini merupakan hal yang umum pada data teks media sosial karena penggunaan kata yang serupa dapat muncul pada konteks yang berbeda.

Meskipun terdapat area yang saling berdekatan, beberapa kluster menunjukkan kecenderungan pemisahan yang lebih baik. Kluster dengan karakteristik bahasa yang lebih spesifik, seperti ekspresi objektifikasi seksual eksplisit, cenderung memiliki pola distribusi yang lebih berbeda dibandingkan kluster lainnya. Sementara itu, kluster yang berkaitan dengan interaksi personal, pujian fisik, atau diskursus gender memiliki distribusi yang lebih berdekatan karena terdapat kemiripan kosakata dan konteks. Hasil PCA ini mendukung interpretasi bahwa komentar misogini memiliki variasi pola bahasa yang kompleks, sehingga pengelompokan tidak sepenuhnya terpisah secara linier.

Hasil dari $k=5$ bisa dilihat pada tabel di bawah ini.

Tabel 2. Top Terms $k=5$

Kluster	Jumlah pernyataan	Top Terms	Contoh	Karakteristik
0	2972	Wanita, susu, gede	Objektifikasi komentar vulgar	Kluster misogini paling jelas
1	415	Perempuan, laki, Indonesia, peran	Diskursus kesetaraan, perempuan	Kluster non misogini / reflektif
2	257	Bagus, cantik, pintar, perempuan	Pujian tapi disertai stereotip	Ambiguitas, pujian bersyarat
3	206	Kamu, aku	Interaksi personal, rayuan	Kluster ambigu, tidak selalu misogini
4	133	Cantik, pacar, orang	Memuji fisik, tetap objektifitas	Sangat relevan dengan literatur gender

Tabel di bawah adalah hasil dari interpretasi metrik clustering dimana untuk menentukan apakah cluster tersebut baik, kurang baik atau tidak baik.

Tabel 3. Interpretasi metrik clustering

Metode	Score	Keterangan
Silhouette Score	0.0195	Nilah rendah dikarenakan kluster tidak terpisah tajam (ciri khas data social)
Davies–Bouldin Index	5.6038	Nilai sedang
Calinski–Harabasz Index	29.5849	Nilai bagus, mendukung bahwa struktur kluster memang ada

KESIMPULAN DAN SARAN

Dengan menggunakan algoritma K-Means, penelitian ini berhasil mengelompokkan komentar media sosial berbahasa Indonesia yang terkait dengan misogini berdasarkan kemiripan semantik. Jumlah kluster yang ideal adalah lima kluster, menurut evaluasi yang dilakukan menggunakan Silhouette Score, Davies-Bouldin Index, dan Calinski-Harabasz Index. Meskipun nilai Silhouette Score agak rendah, kondisi ini masih dapat diterima karena karakteristik data teks media sosial cenderung memiliki tumpang tindih makna dan konteks. Hasil analisis top terms menunjukkan bahwa setiap kluster memiliki karakteristik yang berbeda, mulai dari objektifikasi seksual eksplisit, diskursus gender, benevolent sexism, interaksi interpersonal, hingga objektifikasi fisik yang dinormalisasi. Temuan ini mengindikasikan bahwa ekspresi misogini dalam media sosial memiliki bentuk yang beragam dan tidak dapat dipandang sebagai fenomena yang homogen. Dengan demikian, pendekatan clustering mampu mengungkap struktur semantik laten yang terdapat pada komentar media sosial dan memberikan pemahaman yang lebih mendalam mengenai variasi bentuk misogini dalam bahasa Indonesia.

Untuk penelitian selanjutnya, penggunaan metode ekstraksi fitur berbasis *Deep Learning* seperti *Word Embedding* atau IndoBERT dapat dikaji untuk menangkap hubungan semantik yang lebih kompleks pada komentar media sosial. Selain itu, evaluasi dengan metode clustering lain dan validasi berbasis pakar dapat dilakukan untuk meningkatkan keakuratan interpretasi pola misogini.

DAFTAR PUSTAKA

- [1] Lubis, F. Seksisme dan Misogini dalam Perspektif HAM. Komnas HAM. <https://www.komnasham.go.id/index.php/news/2021/10/28/1963/seksisme-dan-misogini-dalam-perspektif-am.html>, Oct. 28, 2021.
- [2] Maharani, A. P. Misogini: Kebencian Ekstrim Hanya Karena Kita Lembaga Pers Mahasiswa Gelora Sriwijaya. <https://gelorasriwijaya.co/blog/misogini-kebencian-ekstrim-hanya-karena-kita-perempuan/>, Nov. 10, 2022.
- [3] National Democratic Institute. (2019). Tweets That Chill: Analyzing Online Violence Against Women in Politics. <https://www.ndi.org/tweets-that-chill>.
- [4] Damanik, A. J. (2021). IDENTIFIKASI PERNYATAAN MISOGINI BERBAHASA INDONESIA BERDASARKAN KOMENTAR YOUTUBE MENGGUNAKAN GLOVE EMBEDDING DAN RANDOM FOREST CLASSIFIER. In Universitas Sumatera Utara.
- [5] Devi, M. D., & Saharia, N. (2020). Misogynous Text Classification Using SVM and LSTM. In D. Garg, K. Wong, J. Sarangapani, & S. K. Gupta (Eds.), International Advanced Computing Conference 2020 (pp. 336–348). Springer Singapore. <https://doi.org/10.1007/978-981-16-0401-0>.
- [6] Tri Wibowo, B., Nurjanah, D., & Nurrahmi, H. (2023). Identification of Misogyny on Social Media in Indonesian Using Bidirectional Encoder Representations from Transformers (BERT). 5th International Conference on Artificial Intelligence in Information and Communication, ICAIIC <https://doi.org/10.1109/ICAIIIC57133.2023.10067106>.
- [7] Widyadhana, Dahayu, Rina Budi Hastuti, Iqbal Kharisudin, and Fatkhurokhman Fauzi. 2021. “Perbandingan Analisis Kluster K-Means Dan Average Linkage Untuk Pengklasteran Kemiskinan Di Provinsi Jawa Tengah.” PRISMA, Prosiding Seminar Nasional Matematika4(2):584–94.

- [8] Bahauddin, Achmad, Agustina Fatmawati, and Febrianti Permata Sari. 2021. "Analisis Clustering Provinsi Di Indonesia Berdasarkan Tingkat Kemiskinan Menggunakan Algoritma K-Means." *Jurnal Manajemen Informatika Dan Sistem Informasi* 4(1):1–8. doi: 10.36595/misi.v4i1.216.
- [9] afitri, Ayu, Rizki Rahmawati, and Sri Ayu Wandira. 2024. "Pengelompokan Data Penduduk Miskin Di Sumatera Utara Menggunakan K-Means." *J-Com (Journal of Computer)* ISSN 2798-8791 (Online), 20252884(1):33–39. doi: 10.33330/j-com.v4i1.2998.
- [10] Shihab, S.H., Afroge, S., & Mishu, S.Z. 2019. RFM Based Market Segmentation Approach Using Advanced K-Means and Agglomerative Clustering: A Comparative Study. *International Conference on Electrical, Computer and Communication Engineering (ECCE)*, pp. 5386-9111.
- [11] Anggraeni, Lisa, and Priska Arum R. 2022. "Analisis Cluster Menggunakan Algoritma K-Means Pada Provinsi Sumatera Barat Berdasarkan Indeks Pembangunan Manusia Tahun 2021." *Prosiding Seminar Nasional UNIMUS* 5(1):636–46.
- [12] Merriam-Webster. (n.d.). Misogyny. Merriam-Webster.Com Dictionary. Retrieved February 16, 2024, from <https://www.merriam-webster.com/dictionary/misogyny>.