

Pengembangan *Chatbot* Untuk Layanan Informasi Keanggotaan Guru Metode *Support Vector Machine*

Nurkholis Setiawan^{*1}, Muhamad Akbar², Andri Anto Tri Susilo³

Program Studi Informatika, Fakultas Ilmu Teknik, Universitas Bina Insan

Jln. Jendral Besar Moh. Soeharto KM. 13, Kota Lubuk Linggau, Sumatera Selatan

Email: ¹2102020032@mhs.univbinainsan.ac.id, ²muhamad.akbar@univbinainsan.ac.id,

³andri.anto@univbinainsan.ac.id

Abstract. *Development of a Chatbot for Teacher Membership Information Services Using the Support Vector Machine Method.* The growth of chat applications like WhatsApp and Telegram drives the development of chatbots to support intelligent interactions between humans and computers. AI-based chatbots help reduce response times and provide access to information anytime. One issue —repetitive questions about teacher membership administration within the PGRI organization of Musi Rawas Regency —can be addressed with an AI-based chatbot. This study aims to develop a chatbot that delivers information and provides accurate responses. The chatbot is designed using the Support Vector Machine (SVM) algorithm and Natural Language Processing (NLP). Test results show the SVM model achieved 83.54% accuracy, 86.09% precision, and 83.54% recall, making it suitable as a chatbot model. Implementation on the WhatsApp platform makes information easier to access. The results show that the chatbot can respond appropriately to user questions, and integration with WhatsApp makes information services more accessible while providing an efficient, responsive solution.

Keywords: chatbot, Natural Language Processing, Support Vector Machine, teacher membership

Abstrak. *Perkembangan aplikasi chat seperti WhatsApp dan Telegram mendorong pengembangan chatbot untuk mendukung interaksi cerdas antara manusia dan komputer. Chatbot berbasis kecerdasan buatan membantu mengurangi waktu respons dan memberikan akses informasi kapan saja. Salah satu permasalahan, yaitu pertanyaan berulang terkait administrasi keanggotaan guru di PGRI Kabupaten Musi Rawas, dapat diatasi dengan chatbot berbasis AI. Penelitian ini bertujuan mengembangkan chatbot yang mampu menyampaikan informasi dan memberikan respons akurat. Chatbot dirancang dengan algoritma Support Vector Machine (SVM) dan Natural Language Processing (NLP). Hasil pengujian menunjukkan model SVM memiliki akurasi 83,54%, presisi 86,09%, dan recall 83,54%, sehingga layak digunakan sebagai model chatbot. Implementasi pada platform WhatsApp mendukung kemudahan akses informasi. Hasil penelitian membuktikan chatbot dapat merespons pertanyaan sesuai kebutuhan pengguna, sedangkan integrasi dengan WhatsApp membuat layanan informasi lebih mudah diakses serta memberikan solusi yang efisien dan responsif.*

Kata Kunci: chatbot, Natural Language Processing, Support Vector Machine, keanggotaan guru

1. Pendahuluan

Perkembangan teknologi informasi telah merevolusi komunikasi, terutama melalui aplikasi chat seperti WhatsApp, Facebook Messenger, dan Telegram. Awalnya, program-program ini digunakan untuk interaksi manusia ke manusia, tetapi sekarang telah berkembang dengan hadirnya *chatbot* yang memungkinkan percakapan cerdas manusia dan komputer melalui teks atau suara [1]. *Chatbot* telah berkembang tidak hanya menjawab pertanyaan sederhana, tetapi mampu melakukan tugas-tugas administratif dan memberikan layanan khusus. Hal ini membuka peluang penggunaan *chatbot* di organisasi dan lembaga publik untuk meningkatkan kualitas layanan [2].

Salah satunya adalah organisasi Persatuan Guru Republik Indonesia (PGRI) di Kabupaten Musi Rawas, yang memiliki aplikasi KTA Digital sebagai informasi kartu tanda anggota (KTA) bagi setiap guru di wilayahnya. Namun, dalam pengelolaan informasi keanggotaan memberikan informasi kepada anggota dengan cepat dan akurat masih menjadi tantangan. Banyak guru di Kabupaten Musi Rawas yang masih kesulitan memahami proses administrasi yang perlu dilakukan. Pertanyaan yang sama terkait prosedur administrasi umum sering diajukan, sehingga admin harus memberikan penjelasan berulang. Hal ini mengurangi efisiensi waktu dan sedikit menguras tenaga kerja admin. Oleh karena itu diperlukan layanan yang dapat menjawab pertanyaan secara otomatis, seperti *chatbot*.

Chatbot mampu menjawab pertanyaan dengan cepat, konsisten, dan dapat diakses kapan saja tanpa harus menunggu respons dari staf [3]. Beberapa penelitian sebelumnya menunjukkan penggunaan *chatbot* efektif dalam meningkatkan layanan informasi. Salah satu penelitian yang dilakukan oleh D. Akbar et al menggunakan teknologi kecerdasan buatan *Natural Language Processing (NLP)* untuk mengubah layanan pelanggan manual menjadi *chatbot*, hasil penelitian menunjukkan *chatbot* dapat memberikan layanan pelanggan yang lebih cepat, efisien, dan tersedia kapan saja [4]. Berikutnya penelitian oleh M. Mustaqim et al menunjukkan *chatbot* berbasis *machine learning* dan *Natural Language Processing (NLP)* dapat meningkatkan layanan publik dengan lebih cepat, akurat, dan efektif [5].

Machine learning menjadi salah satu pendekatan yang efektif dalam pengembangan *chatbot*. Oleh karena itu, penelitian ini akan mengembangkan *chatbot* berbasis *machine learning* dengan algoritma *Support Vector Machine (SVM)*. Untuk memisahkan berbagai kelas dari data, algoritma *SVM* memiliki kemampuan untuk menentukan *hyperplane* yang ideal. Selain itu, algoritma ini memiliki kinerja yang baik dalam klasifikasi multi-kelas [6]. Seperti pada penelitian D. Lestari & L. Subekti menggunakan algoritma *SVM* untuk klasifikasi *intent* pada *chatbot Telegram*, di mana *SVM* mencapai akurasi sebesar 98,92% melebihi *Naïve Bayes* dan *Neural Network* [7]. Selain itu, penerapan metode *SVM* akan menggunakan pendekatan *One-against-All* untuk menangani masalah klasifikasi *multiclass* pada algoritma *SVM* [8]. Pada penelitian ini, pendekatan *multiclass SVM* bertujuan untuk mengklasifikasikan pola pertanyaan bervariasi namun berulang. Solusi permasalahan tersebut, *chatbot* yang dikembangkan selain *multiclass SVM* selanjutnya akan diimplementasikan melalui *platform WhatsApp* untuk mempermudah penyampaian informasi.

Selanjutnya, penelitian ini difokuskan pada pengembangan *chatbot* yang mampu memberikan layanan informasi terkait proses administrasi pengelolaan keanggotaan PGRI di Kabupaten Musi Rawas dengan respons yang akurat. Dengan menerapkan teknologi kecerdasan buatan dan *platform WhatsApp*, *chatbot* diharapkan dapat memberikan solusi efektif terhadap permasalahan yang dihadapi organisasi PGRI di Kabupaten Musi Rawas dalam mengelola informasi serta meningkatkan kualitas layanan.

2. Tinjauan Pustaka

2.1. Chatbot

Chatbot merupakan sistem berbasis kecerdasan buatan yang dirancang untuk berinteraksi dengan pengguna melalui percakapan teks maupun suara. Teknologi ini banyak digunakan pada berbagai bidang seperti layanan pelanggan, pendidikan, dan layanan publik karena mampu memberikan informasi secara cepat, efisien, dan *real-time* [9] [10]. Dalam konteks organisasi, *chatbot* mampu menggantikan sebagian tugas layanan manual yang sering kali memakan waktu dan sumber daya manusia. Misalnya, penelitian Akbar et al. membuktikan bahwa transformasi layanan *customer service* manual menjadi *chatbot* mampu meningkatkan kecepatan respons dan mengurangi beban kerja petugas [4]. Hal serupa juga ditemukan oleh Mustaqim et al dalam pengembangan *chatbot* layanan publik, di mana penggunaan *chatbot* berbasis *machine learning* dapat meningkatkan kepuasan masyarakat [5]. Untuk organisasi guru seperti PGRI, keberadaan *chatbot* menjadi relevan karena sering kali muncul pertanyaan yang berulang terkait administrasi keanggotaan. Penelitian terdahulu lebih banyak berfokus pada *chatbot* untuk layanan komersial

dan akademik, sehingga penelitian ini berupaya mengisi celah pada ranah layanan informasi keanggotaan guru.

2.2. Natural Language Processing

Natural Language Processing (NLP) adalah bidang ilmu komputer yang mempelajari interaksi antara bahasa manusia dan komputer. *NLP* memungkinkan *chatbot* memahami maksud (*intent*) dari pertanyaan pengguna serta memberikan jawaban yang sesuai [11]. Menurut Malvin & Rangkuti penerapan *NLP* pada *chatbot*, khususnya dalam platform *WhatsApp*, memungkinkan sistem mampu memahami variasi kalimat dan ejaan yang berbeda dari pengguna. Hal ini penting karena dalam layanan publik, gaya bahasa dan struktur kalimat yang digunakan masyarakat sangat bervariasi [12]. *NLP* juga mendukung proses klasifikasi teks pada *chatbot*. Seperti ditunjukkan oleh Putra et al, klasifikasi pesan teks menggunakan *SVM* dalam *chatbot* layanan akademik menghasilkan peningkatan akurasi respons [8]. Dengan demikian, penggunaan *NLP* pada penelitian ini diharapkan dapat meningkatkan keandalan *chatbot* dalam memahami pertanyaan anggota guru yang beragam, khususnya yang berkaitan dengan administrasi keanggotaan.

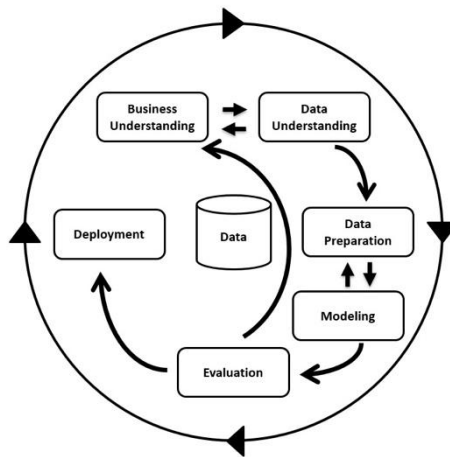
2.3. Support Vector Machine

Support Vector Machine (SVM) merupakan algoritma *supervised learning* yang banyak digunakan dalam klasifikasi teks karena kemampuannya dalam menangani data berdimensi tinggi dan menghasilkan akurasi yang baik [13]. *SVM* bekerja dengan mencari *hyperplane* terbaik yang memisahkan data ke dalam kelas tertentu, sehingga cocok untuk klasifikasi pertanyaan pada *chatbot*. Beberapa penelitian sebelumnya menegaskan keunggulan *SVM*. Habibie Rahman et al. berhasil mengklasifikasikan ujaran kebencian di media sosial dengan tingkat akurasi tinggi menggunakan *SVM* [14]. Fitri & Damayanti juga menunjukkan bahwa *SVM* dapat bersaing dengan algoritma lain seperti *Random Forest* dalam analisis sentimen [6]. Selain itu, Panjaitan et al. membuktikan efektivitas *SVM* dalam mendeteksi spam SMS, yang relevan dengan konteks pertanyaan singkat pada *chatbot* [15].

Khusus pada pengembangan *chatbot*, Lestari & Subekti mengimplementasikan *SVM* untuk klasifikasi teks pada *chatbot Telegram*, yang menunjukkan bahwa *SVM* dapat diandalkan dalam sistem percakapan otomatis. Oleh karena itu, penelitian ini memanfaatkan *SVM* untuk mengklasifikasikan pertanyaan seputar administrasi keanggotaan guru agar *chatbot* mampu memberikan jawaban yang tepat dan konsisten, meskipun pertanyaan diajukan dengan redaksi yang berbeda. Pada penelitian ini untuk menangani masalah klasifikasi multi-kelas, pendekatan *One Against All* diterapkan pada algoritma *SVM*. Pendekatan *One Against All* mengklasifikasikan data ke dalam beberapa kelas dengan membangun sejumlah model *SVM* biner. Pada tahap prediksi, data uji dievaluasi oleh semua model, lalu hasil dengan margin tertinggi dipilih sebagai kelas akhir [14]. Selain itu, dibandingkan dengan pendekatan *One Versus One*, metode *One Against All* dinilai lebih efisien karena menghasilkan performa serupa tetapi dengan waktu pelatihan yang lebih singkat [8]. Sehingga pemilihan pendekatan *One Against All* dilakukan karena data yang digunakan memiliki lebih dari dua kelas *intent*, sehingga metode ini sesuai untuk mendukung proses klasifikasi pertanyaan.

3. Metodologi Penelitian

Penelitian ini menggunakan metode *CRISP-DM (Cross-Industry Standard Process for Data mining)* adalah standar metode yang dikembangkan oleh IBM [13]. Alur tahapan *CRISP-DM* dapat dilihat pada Gambar 1.

Gambar 1. Alur Penelitian Metode *CRISP-DM*

3.1 Business Understanding

Tahap *business understanding* bertujuan untuk memahami tujuan, manfaat, dan kebutuhan pengembangan *chatbot*. Dalam penelitian ini, *chatbot* dirancang untuk memberikan layanan informasi tentang administrasi keanggotaan kepada guru, khususnya menjawab pertanyaan yang sering berulang-ulang tetapi tidak mencakup seluruh varian pertanyaan serupa. Berdasarkan wawancara dengan pengurus organisasi, pertanyaan umum mencakup cara pendaftaran anggota, syarat, dan biaya administrasi.

Tabel 1. Tabulasi distribusi *dataset* awal

No	Tag	Dataset Awal	Dataset Akhir
1	tanya_kta	22	150
2	tanya_daftar	24	150
3	tanya_biaya_admin	24	150
4	tanya_syarat	22	150
5	tanya_status_anggota	22	150
6	tanya_iuran	21	150
7	tanya_verifikasi	24	150
8	sapaan	5	39
9	definisi_ktadigitalpgri	15	33
10	akhiran	5	32
Total		184	1154

Data penelitian diperoleh dari dua sumber. Pertama, hasil wawancara dan observasi yang menghasilkan kumpulan pertanyaan terkait administrasi keanggotaan. Pertanyaan tersebut kemudian diberi label secara manual sesuai *intent* masing-masing dan disimpan dalam format CSV sebagai *dataset* awal. Tabulasi distribusi *dataset* awal yang diperoleh dari hasil pengumpulan data primer disajikan pada Tabel 1. Karena keterbatasan *dataset* yang diperoleh, maka dilakukan pengumpulan *dataset* sumber kedua yaitu melalui proses *augmentasi* data, di mana *dataset* awal diperluas menggunakan teknik parafrase dengan bantuan *ChatGPT-4o* agar memiliki variasi pertanyaan yang lebih beragam [15]. Hasil akhir dari proses ini menghasilkan 1.154 data, sebagaimana ditunjukkan pada Tabel 1 pada kolom 'Dataset Akhir'. Selanjutnya, pada Tabel 2 adalah sampel hasil pengumpulan data.

Tabel 2. Hasil pengumpulan data

Index	Pertanyaan	Tag
0	Bagaimana langkah-langkah untuk cetak kartu PGRI apa bisa dilakukan sendiri? Bagaimana caranya??	tanya_kta
1	Bisakah guru-guru yang ada di sekolah kami mendaftar secara mandiri?	tanya_daftar
2	Apakah guru honor tetap dikenakan biaya admin untuk pendaftaran anggota PGRI?	tanya_biaya_admin

Index	Pertanyaan	Tag
3	Tolong berikan informasi tentang berkas apa saja pak yang dibutuhkan untuk mendaftar anggota PGRI?.	tanya_syarat
4	Bagaimana jika saya ingin tahu kartu anggota PGRI saya sudah aktif atau belum? ? terima kasih.	tanya_status_anggota
5	Apa itu KTA Digital PGRI?	definisi_ktadigitalpgri
6	terimakasih	Akhiran
7	Halo!	Sapaan
...	...	...
1.152	Berapa besar iuran setiap bulan untuk anggota?	tanya_iuran
1.153	Bagaimana cara saya melakukan proses verifikasi anggota PGRI?	tanya_verifikasi

3.2 Data Understanding

Pada tahap *data understanding*, peneliti melakukan eksplorasi dan pemahaman terhadap data yang telah dikumpulkan. Data yang relevan dengan kebutuhan *chatbot* dianalisis dan diproses, kemudian disusun dalam bentuk *dataset* yang terstruktur. Pengumpulan yang dilakukan dengan teknik augmentasi data memerlukan proses penyesuaian ulang secara manual agar data tetap konsisten dan relevan. Selanjutnya, memastikan data dikelompokkan, dikategorikan sesuai konteks pertanyaan, serta penyesuaian label *intent* seperti pendaftaran, syarat, iuran, dan lainnya. Proses ini penting agar setiap data memiliki kelas yang jelas sebelum masuk ke tahap *preprocessing* dan pemodelan.

3.3 Data Preparation

Pada tahap *data preparation*, data yang telah dipahami dan dieksplorasi akan dipersiapkan untuk *training* model *machine learning*. Tahapan dibagi menjadi dua yaitu tahapan pemrosesan teks (*text preprocessing*) dan ekstraksi fitur dengan *TF-IDF*. Tahapan pertama yaitu *Text Preprocessing* yang terdiri dari beberapa langkah penting. Langkah pertama yaitu *case folding*, bertujuan membuat semua huruf dalam teks ke bentuk huruf kecil (*lowercasing*), hal ini bertujuan untuk meningkatkan konsistensi data. Contoh penerapan *case folding* ditunjukkan pada Tabel 3.

Tabel 3. Penerapan *case folding*

Pertanyaan sebelum	Setelah <i>case folding</i>	Tag
Bisakah guru-guru yang ada di sekolah kami mendaftar secara mandiri?	bisakah guru-guru yang ada di sekolah kami mendaftar secara mandiri?	tanya_daftar

Langkah kedua *Remove punctuation* yaitu menghapus tanda baca untuk memastikan teks hanya terdiri dari huruf, angka, atau spasi. Contoh penerapan dari *remove punctuation* ditunjukkan pada Tabel 4.

Tabel 4. Penerapan *Remove Punctuation*

Pertanyaan sebelum	Setelah <i>remove punctuation</i>	Tag
bisakah guru-guru yang ada di sekolah kami mendaftar secara mandiri?	bisakah guru-guru yang ada di sekolah kami mendaftar secara mandiri	tanya_daftar

Langkah ketiga yaitu *Tokenizing* bertujuan untuk memisahkan kata yang ada di suatu teks sehingga menghasilkan data yang berbentuk sebuah *token* (kata-kata). Contoh penerapan dari *tokenizing* ditunjukkan pada Tabel 5.

Tabel 5. Penerapan *Tokenizing*

Pertanyaan sebelum	Setelah <i>tokenizing</i>	Tag
bisakah guru-guru yang ada di sekolah kami mendaftar secara mandiri	['bisakah', 'guruguru', 'yang', 'ada', 'di', 'sekolah', 'kami', 'mendaftar', 'secara', 'mandiri']	tanya_daftar

Langkah keempat *normalization* yaitu bertujuan untuk menyamakan kata yang tidak konsisten seperti pengulangan huruf lebih dari dua kali di akhir kata, contoh ‘aplikasinyaaa’ menjadi ‘aplikasinya’. Selain itu juga seperti pengulangan kata lebih dari dua kali contoh ‘guruguru’ menjadi ‘guru’. Contoh penerapan dari *normalization* ditunjukkan pada Tabel 6.

Tabel 6. Penerapan Normalization

Pertanyaan sebelum	Setelah <i>normalization</i>	Tag
['bisakah', 'guruguru', 'yang', 'ada', 'di', 'sekolah', 'kami', 'mendaftar', 'secara', 'mandiri']	['bisakah', 'guru', 'yang', 'ada', 'di', 'sekolah', 'kami', 'mendaftar', 'secara', 'mandiri']	tanya_daftar

Langkah kelima *stopword removal* bertujuan untuk menghilangkan kata-kata umum yang tidak memberikan informasi terlalu penting. Contoh penerapan dari *stopword removal* ditunjukkan pada Tabel 7.

Tabel 7. Penerapan Stopword Removal

Pertanyaan sebelum	Setelah <i>stopword removal</i>	Tag
['bisakah', 'guru', 'yang', 'ada', 'di', 'sekolah', 'kami', 'mendaftar', 'secara', 'mandiri']	['guru', 'sekolah', 'mendaftar', 'mandiri']	tanya_daftar

Langkah keenam *stemming* bertujuan mengembalikan kata berimbuhan ke bentuk aslinya sesuai dengan kamus bahasa yang ditentukan, contohnya seperti ‘mendaftar’ menjadi ‘daftar’. Contoh penerapan dari *stemming* ditunjukkan pada Tabel 8.

Tabel 8. Penerapan Stemming

Pertanyaan sebelum	Setelah <i>stemming</i>	Tag
['guru', 'sekolah', 'mendaftar', 'mandiri']	['guru', 'sekolah', 'daftar', 'mandiri']	tanya_daftar

Tahap kedua adalah ekstraksi fitur menggunakan *TF-IDF* (*Term Frequency-Inverse Document Frequency*), yang digunakan untuk menghitung tingkat kepentingan suatu kata dalam dokumen.. Perhitungan terdiri dari tiga yaitu *TF* (*Term Frequency*) yang dinyatakan dalam Persamaan 1 dan *IDF* (*Inverse Document Frequency*) dinyatakan dalam Persamaan 2, kemudian dilakukan perhitungan akhir antara *TF* dan *IDF* yang dinyatakan dalam Persamaan 3 [16].

$$TF(t, d) = \frac{\text{Jumlah kemunculan kata } t \text{ dalam dokumen } d}{\text{Jumlah kata dalam dokumen } d} \quad (1)$$

$$IDF(t, D) = \log\left(\frac{\text{Jumlah total dokumen } N}{\text{Jumlah dokumen yang mengandung kata } t+1}\right) \quad (2)$$

$$TF - IDF(t, d, D) = TF(t, d) \times IDF(t, D) \quad (3)$$

3.4 Modelling

Setelah data dikonversi ke *TF-IDF*, dilakukan pemodelan klasifikasi menggunakan algoritma *SVM* dengan *kernel* linear, karena cocok untuk data berdimensi tinggi seperti *TF-IDF* dan umum digunakan dalam klasifikasi teks [13]. Untuk klasifikasi multi-kelas, digunakan pendekatan *One Against All*, di mana setiap kelas dibandingkan sebagai kelas positif (label 1) dan sisanya sebagai kelas negatif (label -1). Model dilatih menggunakan *dataset* yang telah dibagi menjadi data latih dan data uji [14].

3.5 Evaluation

Tahap *evaluation* bertujuan mengukur performa model *chatbot* setelah dilakukan pelatihan. Evaluasi dilakukan menggunakan *Confusion Matrix* dan *Cross-Validation*, yang

selanjutnya digunakan untuk menghitung metrik seperti nilai akurasi pada Persamaan 4, nilai Presisi pada Persamaan 5, dan nilai *Recall* pada Persamaan 6.

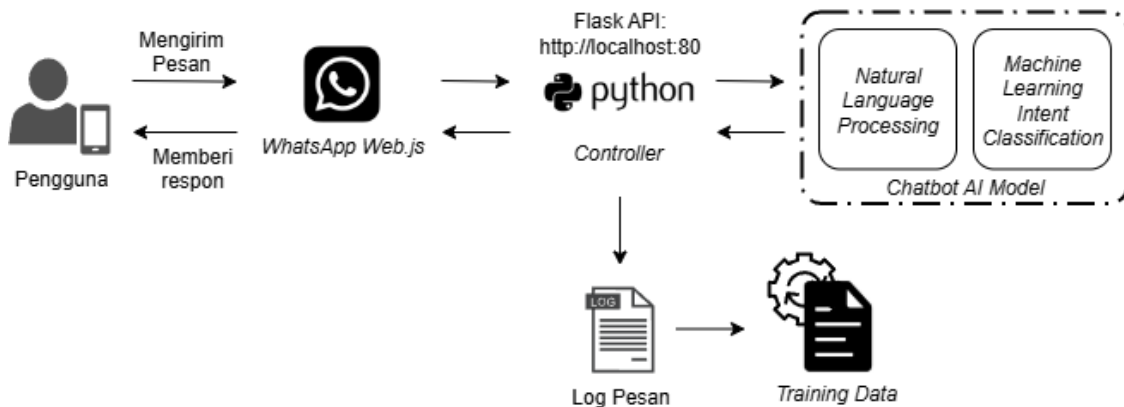
$$\text{Akurasi} = \frac{TP + TN}{TP + TN + FP + FN} \quad (4)$$

$$\text{Presisi} = \frac{TP}{TP + FP} \quad (5)$$

$$\text{Recall} = \frac{TP}{TP + FN} \quad (6)$$

3.6 Deployment

Terakhir, pada tahapan *deployment*, hasil model yang telah diuji dan siap digunakan akan diimplementasikan untuk melakukan klasifikasi *intent* (maksud) pertanyaan pengguna secara *real-time*. Model ini akan diterapkan di *backend server* secara lokal dengan memanfaatkan framework *Flask Python* sebagai *application programming interface (API)*, yang kemudian diintegrasikan dengan platform *WhatsApp* untuk menyediakan layanan interaktif kepada pengguna. Rancangan arsitektur *deployment chatbot* disajikan pada Gambar 2.



Gambar 2. Rancangan Arsitektur *Deployment Chatbot*

4. Hasil dan Diskusi

Pada penelitian ini dilakukan dua tahap pengujian. Pertama, menggunakan data uji dari *dataset* campuran yang terdiri atas data primer dan hasil augmentasi yang telah digabungkan serta melalui tahap *preprocessing* hingga pembagian data latih 75% dan data uji 25%. Kedua, menggunakan data uji primer murni yang berasal dari pertanyaan baru dan tidak termasuk ke dalam *dataset* campuran dengan augmentasi. Hal ini bertujuan untuk mengidentifikasi potensi terjadinya bias akibat penggunaan *dataset* augmentasi, sekaligus memastikan bahwa model tetap mampu mengenali pola pertanyaan. Berikut adalah komposisi data uji masing-masing kelas ditunjukkan pada Tabel 9.

Tabel 9. Distribusi *Dataset Uji*

No	Tag	Primer & Augmentasi	Primer
1	tanya kta	31	20
2	tanya daftar	38	20
3	tanya biaya admin	41	20
4	tanya syarat	36	20
5	tanya status anggota	42	20
6	tanya iuran	44	20
7	tanya verifikasi	36	20
8	sapaan	9	6
9	definisi ktadigitalpgri	6	6
10	akhiran	6	6
Total		289	158

Jumlah data uji primer baru ditentukan secara proporsional dan seimbang antar kelas tujuannya adalah untuk menghindari ketidakseimbangan distribusi data uji sehingga evaluasi performa model lebih adil. Pendekatan ini dilakukan agar potensi bias akibat data augmentasi dapat diminimalkan dan kemampuan model dalam mengenali pertanyaan pengguna dapat terukur dengan baik.

4.1 Hasil Evaluasi Model

Evaluasi model *Support Vector Machine* pertama dilakukan dengan menggunakan *confusion matrix* untuk menghitung nilai metrik *akurasi*, *presisi*, dan *recall*. Pada penelitian ini nilai akurasi pada data uji primer tanpa augmentasi data, yang akan diambil sebagai kesimpulan hasil penelitian. Hasil evaluasi model *SVM* masing-masing data uji sebagaimana ditunjukkan pada Tabel 10.

Tabel 10. Hasil Uji Performa Model SVM

<i>Dataset Uji</i>	<i>Akurasi</i>	<i>Presisi</i>	<i>Recall</i>
Augmentasi dan Primer	95,89	96,95	95,89
Primer	83,54	86,09	83,54

Hasil pengujian memperlihatkan adanya perbedaan signifikan antara *dataset* campuran (primer dan augmentasi) dengan *dataset* primer baru. Pada *dataset* campuran, model memperoleh akurasi 95,89% dengan presisi 96,95% dan *recall* 95,89%. Hal ini menunjukkan bahwa penggunaan data augmentasi mampu menambah variasi pertanyaan sehingga model lebih akurat dan stabil dalam melakukan klasifikasi, meskipun dapat menimbulkan potensi bias karena pola pertanyaan sama dalam proses augmentasi *dataset*.

Sementara itu, pada *dataset* primer, performa menurun dengan akurasi 83,54%, presisi 86,09%, dan *recall* 83,54%. Perbedaan sekitar 12% pada akurasi ini menandakan bahwa tanpa augmentasi, model lebih rentan terhadap keterbatasan variasi pertanyaan baru. Hasil ini membuktikan bahwa meskipun performa sedikit menurun, model tetap mampu mengenali pola pertanyaan non-augmentasi dengan cukup baik. Dengan demikian, penggunaan data augmentasi memang meningkatkan keragaman *dataset*, tetapi model tidak sepenuhnya bergantung pada pola hasil augmentasi. Hal ini menegaskan bahwa model tidak bias terhadap data augmentasi dan masih mampu digunakan untuk mengenali pertanyaan baru dengan baik. Untuk memastikan konsistensi performa model, dilakukan evaluasi menggunakan *10-fold cross validation* pada data uji primer baru. Hasil pengujian *10-fold cross validation* ditunjukkan pada Tabel 11.

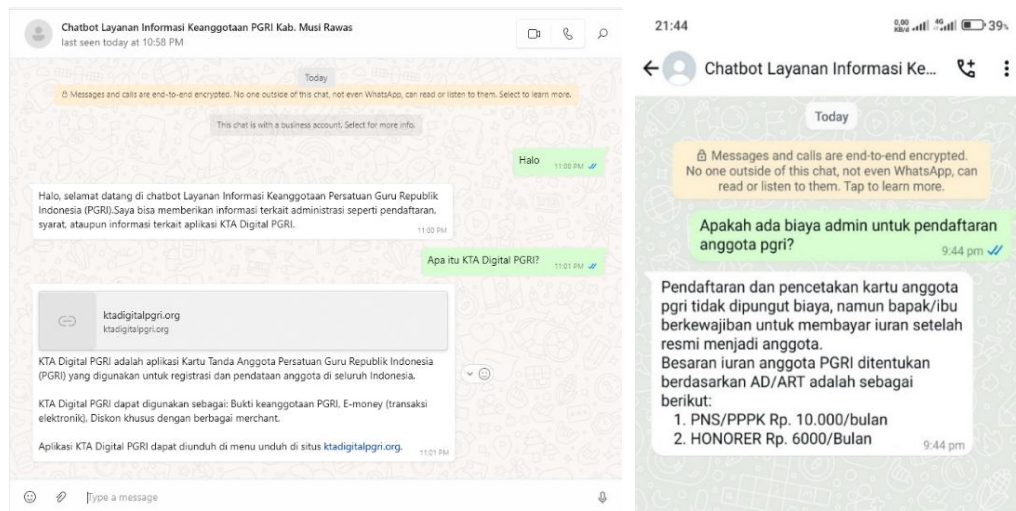
Tabel 11. Hasil K-Fold Cross Validation

<i>K-Fold</i>	<i>Score</i>
Fold 1	0,941176
Fold 2	0,823529
Fold 3	0,941176
Fold 4	0,823529
Fold 5	0,812500
Fold 6	0,875000
Fold 7	0,875000
Fold 8	0,937500
Fold 9	0,812500
Fold 10	0,875000
Rata-rata	0,871691

Evaluasi menggunakan *10-fold cross validation* pada data primer menghasilkan skor rata-rata sebesar 0,87 dengan variasi skor antara 0,81 hingga 0,94 di setiap *fold*. Hasil ini menunjukkan bahwa model memiliki performa yang stabil pada data primer baru dan mampu mempertahankan konsistensi klasifikasi meskipun tanpa dukungan data augmentasi. Dengan demikian, *cross validation* memperkuat bukti bahwa model tetap reliabel dan tidak bias, serta layak diterapkan untuk mengenali pertanyaan dari pengguna.

4.2 Hasil Deployment

Model *TF-IDF* dan *SVM* yang telah dikembangkan melalui berbagai tahap selanjutnya diimplementasikan pada sistem *chatbot*. *Chatbot* dirancang menggunakan *library whatsapp-web.js* sebagai interaksi pesan *chatbot* dengan pengguna melalui aplikasi *WhatsApp*. Demo hasil *deployment* dapat dilihat pada Gambar 3.



Gambar 3. Demo Chatbot pada platform WhatsApp

Dapat dilihat pada Gambar 3 hasil *deployment* model ke platform *WhatsApp*, di mana pengguna memberikan pertanyaan dan *chatbot* memberikan respons yang sesuai. Proses ini dilakukan melalui klasifikasi *intent*, kemudian *intent* tersebut memanggil data jawaban yang telah disiapkan. Setiap data jawaban dapat mewakili banyak pertanyaan yang masih berada dalam satu konteks, sehingga meskipun variasi pertanyaan berbeda, respons yang diberikan bisa tetap relevan.

5. Kesimpulan dan Saran

Berdasarkan hasil penelitian, model *Support Vector Machine (SVM)* pada pengujian data uji primer memiliki nilai akurasi 83,54%, presisi 86,09%, dan *recall* 83,54%. Hasil ini menunjukkan bahwa pengujian tanpa data augmentasi model tetap mampu mengenali pola pertanyaan dengan cukup baik. Evaluasi lebih lanjut menggunakan *10-fold cross validation* menghasilkan rata-rata skor 0,87 dengan variasi skor relatif kecil, yang menegaskan bahwa performa model stabil, konsisten, dan tidak bias terhadap pola tertentu dari data pelatihan. Lebih lanjut Integrasi *chatbot* dengan *WhatsApp* dapat meningkatkan aksesibilitas layanan informasi, karena pengguna dapat mengaksesnya tanpa aplikasi tambahan. Hal ini membuat layanan lebih efisien, cepat, dan responsif bagi anggota PGRI di Kabupaten Musi Rawas. Saran penelitian selanjutnya pengembangan *chatbot* tanpa harus terbatas pada pendekatan klasifikasi tradisional. Sebagai alternatif, dapat digunakan model canggih yang sudah tersedia seperti *pretrained* model atau *deep learning*, sehingga *chatbot* mampu memberikan respons yang lebih akurat, alami, dan sesuai dengan konteks pertanyaan pengguna.

Referensi

- [1] D. A. Dzulhijjah and R. Darodjat, "Analisis Kebutuhan Fungsional Sistem Dalam Merancang Chatbot Tiket Kapal Pt Xyz Dengan Metode Pieces," *Prosiding Seminar Nasional Sains dan Teknologi (SAINTEK)*, vol. 1, no. 1, pp. 229–237, 2023.
- [2] M. P. Wicaksana, P. G. Rahardandi, and M. Fauzan, "Analisis Penerapan Chatbot: Survei," *Innovative: Journal Of Social Science Research*, vol. 4, no. 4, pp. 8349–8364, 2024, doi: 10.31004/innovative.v4i4.13789.

- [3] L. Anindyati, "Analisis dan Perancangan Aplikasi Chatbot Menggunakan Framework Rasa dan Sistem Informasi Pemeliharaan Aplikasi (Studi Kasus: Chatbot Penerimaan Mahasiswa Baru Politeknik Astra)," *Jurnal Teknologi Informasi dan Ilmu Komputer (JTIIK)*, vol. 10, no. 2, pp. 291–300, 2023, doi: 10.25126/jtiik.20231026409.
- [4] D. Akbar, I. Rizky Priadi, W. B. Pamungkas, W. Oetama, and A. Saifudin, "Transformasi Sistem Customer Service Manual Menjadi Chatbot Memanfaatkan Kecerdasan Buatan Natural Language Processing Di Pt.Xyz," *BIIKMA : Buletin Ilmiah Ilmu Komputer dan Multimedia*, vol. 2, no. 2, pp. 393–397, 2024.
- [5] M. Mustaqim, A. Gunawan, Y. B. Pratama, and I. Zaliman, "Pengembangan Chatbot Layanan Publik Menggunakan Machine Learning Dan Natural Language Processing," *Journal of Information Technology and society (JITS)*, vol. 1, no. 1, pp. 1–4, 2023, doi: 10.35438/jits.v1i1.16.
- [6] D. A. Fitri and Damayanti, "Komparasi Algoritma Random Forest Classifier dan Support Vector Machine untuk Sentimen Masyarakat terhadap Pinjaman Online di Media Sosial," *JUPI (Jurnal Ilmiah Penelitian dan Pembelajaran Informatika)*, vol. 9, no. 4, pp. 2018–2029, 2024, doi: 10.29100/jupi.v9i4.5608.
- [7] D. Lestari and L. Subekti, "Implementasi Chatbot pada Telegram sebagai Monitoring Assistant dengan Analisis Teks Klasifikasi Menggunakan Metode Support Vector Machine," *Journal of Internet and Software Engineering*, vol. 5, no. 2, pp. 68–74, 2024, doi: 10.22146/jise.v5i2.12928.
- [8] R. P. Putra, A. H. Pratomo, and R. I. Perwira, "Text Message Classification using Multiclass Support Vector Machine on Information Service Chatbot in the Informatics Department UPN 'Veteran' Yogyakarta," *Telematika: Jurnal Informatika dan Teknologi Informasi*, vol. 19, no. 3, pp. 295–310, 2022, doi: 10.31315/telematika.v19i3.7418.
- [9] E. Erlina *et al.*, "Penerapan Artificial Intelligence pada Aplikasi Chatbot sebagai Sistem Pelayanan dan Informasi Online pada Sekolah," *Journal of Information System and Technology*, vol. 4, no. 3, pp. 221–230, 2023, doi: 10.37253/joint.v4i3.6296.
- [10] A. Lubis and I. Sumartono, "Implementasi Layanan Akademik Berbasis Chatbot untuk Meningkatkan Interaksi Mahasiswa," *Media Online*, vol. 3, no. 5, pp. 397–403, 2023, doi: 10.30865/resolusi.v3i5.767.
- [11] M. A. Nadzif, Saefurrohman, and R. Soelistijadi, "Penggunaan Teknologi Natural Language Processing dalam Sistem Chatbot untuk Peningkatan Layanan Informasi Administrasi Publik," *Indonesian Journal of Computer Science*, vol. 13, no. 1, pp. 1227–1242, 2024, doi: 10.33022/ijcs.v13i1.3645.
- [12] Malvin, C. Dylan, and A. H. Rangkuti, "WhatsApp Chatbot Customer Service Using Natural Language Processing and Support Vector Machine," *International Journal of Emerging Technology and Advanced Engineering (IJETA)*, vol. 12, no. 3, pp. 130–136, 2022, doi: 10.46338/ijetae0322_15.
- [13] F. Abdusyukur, "Penerapan Algoritma Support Vector Machine (Svm) Untuk Klasifikasi Pencemaran Nama Baik Di Media Sosial Twitter," *Komputa : Jurnal Ilmiah Komputer dan Informatika*, vol. 12, no. 1, pp. 73–82, 2023, doi: 10.34010/komputa.v12i1.9418.
- [14] O. Habibie Rahman, G. Abdillah, and A. Komarudin, "Klasifikasi Ujaran Kebencian pada Media Sosial Twitter Menggunakan Support Vector Machine," *Rekayasa Sistem dan Teknologi Informasi (RESTI)*, vol. 1, no. 10, pp. 17–23, 2021, doi: 10.29207/resti.v5i1.2700.
- [15] M. D. Panjaitan, P. P. Adikara, and B. D. Setiawan, "Klasifikasi Spam pada Short Message Service (SMS) menggunakan Support Vector Machine," *Jurnal Pengembangan Teknologi Informasi dan Ilmu Komputer (JPTIIK)*, vol. 8, no. 5, pp. 1–10, 2024.
- [16] K. Tri Putra, M. Amin Hariyadi, and C. Crysdiyan, "Perbandingan Feature Extraction Tf-Idf Dan Bow Untuk Analisis Sentimen Berbasis SVM," *Jurnal Cahaya Mandalika (JCM)*, vol. 3, no. 3, pp. 1449–1463, 2023.