

## KLASIFIKASI PENERIMA BANTUAN PROGRAM KELUARGA HARAPAN MENGUNAKAN SUPPORT VECTOR MACHINE

Moch. Nur Isyam, Didik Indrayana, Winda Apriandari

Teknik Informatika, Universitas Muhammadiyah Sukabumi

Jl. R. Syamsudin, S.H. No. 50, Cikole, Kec. Cikole, Kota Sukabumi, Jawa Barat 43113

mnurisyam043@ummi.ac.id

### ABSTRAK

Program Keluarga Harapan adalah sebuah program bantuan sosial dari pemerintah Indonesia yang bertujuan untuk mengurangi kemiskinan melalui bantuan keuangan untuk keluarga kurang mampu. Fokus penelitian ini adalah mengatasi masalah subjektivitas dalam pemilihan penerima bantuan di Desa Cisarua, Kecamatan Nagrak, Kabupaten Sukabumi, proses pemilihan yang masih bergantung pada musyawarah desa dan pendataan secara manual seringkali kurang objektif dan rentan terhadap kesalahan karena banyaknya data yang harus dikelola secara manual. Untuk mengatasi masalah tersebut peneliti menerapkan metode *Support Vector Machine* untuk klasifikasi penerima bantuan. Penggunaan SVM untuk klasifikasi penerima bantuan menghasilkan akurasi sebesar 89,89%, precision 90%, recall 92% dan F1-Score 91%. Ini menunjukkan efektivitasnya dalam klasifikasi penerima bantuan PKH.

**Kata Kunci :** *Support Vector Machine, Program Keluarga Harapan, Klasifikasi*

### 1. PENDAHULUAN

Masalah kemiskinan di Indonesia adalah sebuah masalah mendasar yang penting untuk diatasi diseluruh dunia termasuk juga di Indonesia[1].

Data dari Badan Pusat Statistik Indonesia mencatat bahwa presentase penduduk Indonesia mencapai 9,36% pada maret 2023 [2]. Pemerintah Indonesia berupaya menangani masalah ini secara cepat dan efektif.

Program Keluarga Harapan atau PKH adalah salah satu usaha pemerintah untuk mengurangi kemiskinan di Indonesia. Sejak dibentuk pada tahun 2007, PKH dirancang untuk membantu keluarga yang tidak mampu untuk menjaga daya beli mereka, terutama pada saat terjadinya perubahan harga bahan bakar [3].

Bantuan ini memberikan bantuan berupa tunai bersyarat kepada keluarga miskin (KSM/RSTM) yang telah terdaftar sebagai peserta dengan memenuhi kriteria tertentu [4].

Tujuan utama dari PKH adalah untuk meningkatkan akses dan kualitas pelayanan dibidang pendidikan, dan kesehatan bagi peserta, yang diharapkan dapat berdampak positif pada kesejahteraan masyarakat secara keseluruhan [5].

Di Desa Cisarua, Kecamatan Nagrak, Kabupaten Sukabumi, penentuan keluarga yang layak menerima bantuan PKH tidak memiliki acuan yang jelas, penentuan dilakukan dengan musyawarah dan pendataan secara manual yang seringkali rentan terhadap kesalahan akibat banyaknya data. Hal ini dapat menyebabkan ketidakakuratan dalam pemilihan penerima bantuan. Sedangkan pemerintah daerah menjadi bagian penting untuk kepesertaan PKH dalam suatu daerah seperti yang ada di dalam peraturan Menteri Sosial Republik Indonesia Nomor 1 Tahun 2018 tentang Program Keluarga Harapan pada bab 5 bagian ketiga mengenai Penetapan Calon Peserta PKH

pada pasal 34 disana diterangkan bahwa “Data tingkat kemiskinan dan kesiapan pemerintah daerah menjadi salah satu bahan pertimbangan dalam penetapan wilayah kepesertaan PKH” [6].

Oleh karena itu untuk menentukan penerima bantuan peneliti mencoba untuk membuat sebuah sistem berbasis *machine learning* yang dapat mengklasifikasi dan memberikan acuan penerima bantuan. *Machine learning* sendiri merupakan penerapan komputer dan algoritma matematika untuk menghasilkan sebuah prediksi dari pembelajaran data, terbagi kedalam beberapa kategori yaitu *Supervised Learning*, *Unsupervised Learning*, dan *Reinforcement Learning* [7].

Untuk Klasifikasinya menggunakan *Support vector Machine* dalam kategori *Supervised Learning* untuk mengklasifikasikan penerima bantuan. SVM efektif dalam menganalisis data dan mengenali pola untuk klasifikasi [8].

### 2. TINJAUAN PUSTAKA

#### 2.1. Program Keluarga Harapan

Program Keluarga Harapan adalah sebuah program yang dikeluarkan pemerintah untuk menangani masalah kemiskinan di Indonesia, diperuntukan untuk rumah tangga sangat miskin (RSTM) yang termasuk ke dalam persyaratan yang sudah ditentukan. Selain itu program ini merupakan bentuk upaya peningkatan sumber data manusia (SDM). Bantuan PKH ini dalam konsepnya merupakan program jaminan sosial yang berbentuk tunjangan tunai, barang ataupun pelayanan kesejahteraan [9].

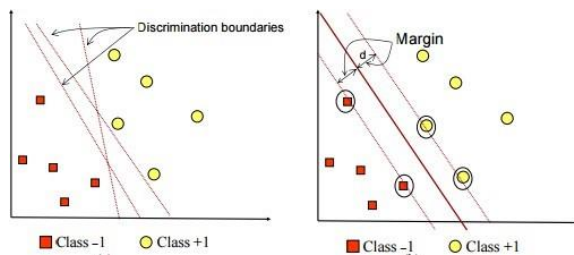
#### 2.2. Machine Learning

*Machine Learning* adalah pembelajaran data untuk menghasilkan prediksi dimasa yang akan datang dengan penerapan komputer dan algoritma matematika.

Pembelajaran dalam konteks ini terbagi ke dalam dua tahap utama yaitu *training* dan *testing*. Terbagi atas tiga kategori yaitu *supervised Learning*, *Unsupervised Learning* dan *Reinforcement Learning* [7]

### 2.3. Support Vector Machine

*Support Vector Machine* merupakan sebuah sistem pembelajaran yang mempelajari pola dengan memahami ruang hipotesis yang berupaya untuk membuat fungsi – fungsi linear dalam sebuah ruang fitur (*feature space*) berdimensi tinggi. Algoritma ini memanfaatkan learning bias dan optimasi untuk mendapatkan hasil yang baik [10].



Gambar 1. Pemisahan Kelas SVM

Ada dua pola yang dikelompokkan dalam dua kelas -1 berwarna merah dan kelas +1 berwarna kuning. Terdapat dua situasi dalam menentukan pemisahan kelas dengan menggunakan hyperlane yaitu kelas yang dapat dipisahkan secara sempurna disebut SVM linier dan kelas yang tidak dapat dipisahkan secara sempurna disebut SVM non-linier [11].

### 2.4. Klasifikasi

Klasifikasi adalah sebuah proses menilai objek data untuk mengelompokkannya ke dalam kelas tertentu dari sejumlah kelas yang berbeda. Proses ini melibatkan pembuatan suatu model berdasar kepada data pelatihan yang ada, kemudian menggunakan model yang sudah dilatih untuk mengklasifikasikan data baru. Definisi klasifikasi bisa disebut sebagai kegiatan yang melibatkan pelatihan atau pembelajaran fungsi target yang memetakan setiap set atribut ke salah satu label yang tersedia [12].

### 2.5. Knowledge Discovery in Database (KDD)

Knowledge Discovery in Database merupakan sebuah tahapan yang meliputi pengumpulan data, pemakaian data, historis untuk menemukan keteraturan pola atau hubungan dalam sebuah dataset yang memiliki ukuran besar ada beberapa tahapan dari KDD yaitu :

#### a. Data Selection

Tahap ini mencakup pemilihan data yang relevan dengan tujuan penelitian. Data yang dipilih harus sesuai dengan kebutuhan penelitian untuk memastikan hasil yang tepat. Setelah memilih data yang sesuai, informasi tersebut disimpan secara terpisah, sehingga memudahkan pengelolaan dan menjaga integritas data selama proses penelitian.

#### b. Preprocessing

Tahap ini sangat penting karena berfokus pada pembersihan data dari masalah umum seperti duplikasi dan inkonsistensi. Selain itu, pre-processing juga mencakup penambahan data relevan untuk memperkaya dataset yang akan dianalisis. Langkah ini memastikan bahwa data yang digunakan dalam analisis lebih lengkap dan mencakup semua aspek yang diperlukan untuk penelitian.

#### c. Transformation

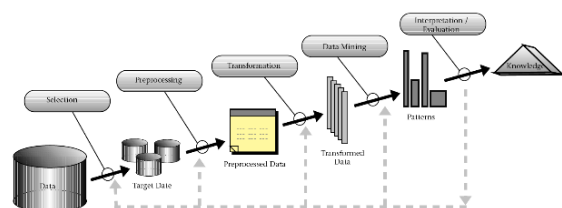
Pada tahap ini, data yang telah dibersihkan dan diperkaya diubah ke dalam format yang lebih sesuai untuk analisis. Proses transformasi ini melibatkan pengkodean atau penyesuaian skala data sesuai dengan pola informasi yang relevan untuk penelitian. Transformasi bertujuan untuk mempermudah identifikasi pola dan hubungan dalam data pada tahap analisis selanjutnya.

#### d. Data Mining

Pada tahap ini, berbagai metode dan algoritma diterapkan untuk mengeksplorasi dan mengidentifikasi pola dalam data. Peneliti sering menggunakan perangkat lunak seperti Google Colab, yang memfasilitasi penerapan berbagai algoritma untuk mengungkap hubungan yang mungkin tidak jelas atau tersembunyi dalam data yang besar. Pemilihan metode tergantung pada jenis data dan tujuan spesifik dari penelitian.

#### e. Evaluation

Ini adalah proses verifikasi di mana pola yang ditemukan selama data mining diuji untuk menentukan keakuratan dan relevansinya terhadap hipotesis atau pertanyaan penelitian. Evaluasi ini penting untuk memastikan bahwa hasil dari data mining dapat diandalkan dan valid. Apabila hasil evaluasi menunjukkan ketidaksesuaian dengan yang diharapkan, mungkin diperlukan penyesuaian pada proses analisis atau mungkin ada kebutuhan untuk kembali ke tahap awal dan memperbaiki pemilihan atau pemrosesan data [13]



Gambar 2. Tahapan KDD

### 2.6. Google Colaboratory

Google Colaboratory merupakan produk dari Google Research yang memungkinkan penggunaanya untuk menulis dan menjalankan kode Python melalui browser. Platform ini sangat cocok untuk keperluan seperti machine learning, analisis data, dan akademik. Secara teknis, Google Colab ini adalah layanan notebook seperti Jupyter yang dihosting dan dapat digunakan tanpa memerlukan penyimpanan lokal [14].

### 2.7. Flask

Flask, sebagai contoh, adalah microframework web yang dibangun dengan bahasa pemrograman Python. Flask dirancang untuk bekerja secara efisien pada perangkat dengan keterbatasan daya dan memori. Meskipun ringan, Flask menyediakan berbagai fungsi yang dapat dikembangkan sesuai dengan kebutuhan pengguna. Selain itu, Flask juga mendukung integrasi yang baik dengan database melalui penggunaan SQLAlchemy [15]

## 3. METODE PENELITIAN

Metode penelitian memiliki tujuan untuk memberikan pandangan yang rinci terkait rencana pelaksanaan penelitian. Selain menguraikan dengan detail proses pengumpulan data berdasarkan informasi dari instansi yang menjadi fokus penelitian, tahap ini juga mencakup penjelasan lebih lanjut mengenai kebutuhan yang diperlukan dalam menyelesaikan penelitian tersebut.

### 3.1. Data Selection

Pada tahap ini data yang didapatkan diolah melalui proses penyeleksian data, Dimana data ini diseleksi berdasarkan kebutuhan pada penelitian ini. Data yang didapatkan dari instansi penelitian yaitu Desa Cisarua Kecamatan Nagrak Kabupaten Sukabumi berupa data penerima pkh yang layak dan tidak layak yang berjumlah 442. dengan data yang berisi nama, nik, alamat, lokasi, usia, pekerjaan, penghasilan, jumlah keluarga, jumlah tanggungan anak, status rumah, jenis lantai, jenis dinding dan Kelayakan. Penyeleksian data atau pemilihan atribut dilakukan dari yang awalnya berjumlah 13 menjadi 9 kolom. Kolom yang tidak digunakan diantaranya nama, nik, alamat dan lokasi.

### 3.2. Preprocessing Data

Dalam tahap ini adalah mengelola data mentah menjadi data yang sesuai dengan data yang akan dianalisis selanjutnya. Preprocessing data ini meliputi antara lain adalah menghapus data duplikasi, memeriksa data yang inkonsisten dan memperbaiki kesalahan pada data. Preprocessing data ini dapat meningkatkan kualitas data serta meningkatkan nilai akurasi dan efisiensi proses data mining.

### 3.3. Tranformation

Pada tahap transformasi data PKH, data dibagi menjadi Data Training dan Data Testing untuk memvalidasi model. Juga dilakukan perubahan variabel ke format yang lebih mudah dianalisis untuk analisis statistik dan model prediktif.

### 3.4. Data Mining

Pembuatan pemodelan dengan algoritma SVM melibatkan penentuan jenis kernel yang paling sesuai dengan karakteristik data dan tujuan analisis kita, dan evaluasi kinerja model pada data testing untuk memastikan akurasi prediksi pada data baru.

### 3.5. Evaluation

Pada tahap ini, model SVM dilatih dengan subset data latih yang telah diproses, menentukan parameter optimal untuk kinerja maksimal. Setelah pelatihan, model dievaluasi dengan data uji untuk mengukur akurasi dan kemampuan prediksi, memastikan generalisasi yang baik untuk data baru.

### 3.6. Deployment

Untuk kemudahan dalam penggunaan peneliti membuat sebuah sistem berbasis web untuk klasifikasi penerima bantuan, dengan menambahkan model yang telah dihasilkan.

## 4. HASIL DAN PEMBAHASAN

### 4.1. Data Selection

Data yang sudah diperoleh atau dikumpulkan mencakup beberapa variabel, yaitu :

Tabel 1. Seleksi Atribut

| Atribut                | Detail Penggunaan |              |
|------------------------|-------------------|--------------|
| No. KK                 | X                 | Nilai Model  |
| Lokasi                 | X                 | Nilai Model  |
| Nama                   | √                 | Nilai Model  |
| Alamat                 | X                 | Nilai Model  |
| Usia                   | √                 | Nilai Model  |
| Pekerjaan              | √                 | Nilai Model  |
| Penghasilan            | √                 | Nilai Model  |
| Jumlah Keluarga        | √                 | Nilai Model  |
| Jumlah Tanggungan Anak | √                 | Nilai Model  |
| Status Rumah           | √                 | Nilai Model  |
| Jenis Lantai           | √                 | Nilai Model  |
| Jenis Dinding          | √                 | Nilai Model  |
| Status                 | √                 | Label Target |

Tabel di atas menunjukkan atribut yang digunakan dalam penelitian dengan tanda ceklis (√) dan atribut yang dihapus dengan tanda silang (×).

Setelah pemilihan atribut, data akan dikonversi untuk mempermudah pemodelan.

### 4.2. Preprocessing Data

Pada tahap ini dilakukan persiapan data diantaranya menghapus atribut yang tidak digunakan, mencari nilai yang hilang dan melakukan perubahan pada atribut ['Penghasilan'] agar mudah diproses saat pemodelan.

```
# Menghapus kolom yang tidak perlu
df = df.drop(columns=['NO KK', 'Alamat', 'Lokasi'])
```

Menghapus kolom kolom yang tidak diperlukan.

```
df.isnull().sum()
Usia      0
Pekerjaan 0
Penghasilan 0
Jumlah Keluarga 0
Jumlah Tanggungan Anak 0
Status Rumah 0
Jenis Lantai 0
Jenis Dinding 0
Status 0
dtype: int64
```

Gambar 3. Missing Value

Missing Value apabila tidak ditangani akan berpengaruh terhadap pemodelan nantinya. Dilakukan pengkategorian untuk data string dan numerik juga dilakukan perubahan bentuk untuk kolom 'Penghasilan' agar lebih mudah diproses.

```
categorical_columns = ['Pekerjaan', 'Status Rumah',
                       'Jenis Lantai', 'Jenis Dinding', 'Penghasilan']
numerical_columns = ['Usia', 'Jumlah Keluarga',
                      'Jumlah Tanggungan Anak']

Def categorize_penghasilan(value):
    value = value.strip()
    if value == '< 1.000.000':
        return 'rendah'
    elif value == '< 2.000.000' or value == '< 3.000.000':
        return 'sedang'
    elif value == '< 4.000.000' or value == '> 4.000.000':
        return 'tinggi'
    else:
        raise ValueError(f'Unexpected value: {value}')
df['Penghasilan'] = df['Penghasilan'].apply(categorize_penghasilan)
```

Cuplikan dari kode diatas adalah identifikasi kolom dilakukan dengan memisahkan kolom string dan numerik. categorical\_columns adalah daftar kolom dengan data string, sedangkan numerical\_columns adalah daftar kolom dengan data numerik. Selanjutnya kolom penghasilan dirubah menjadi kategori yaitu < 1.000.000 dan < 2.000.000 menjadi "rendah", < 3.000.000 sebagai "sedang", dan > 4.000.000 sebagai "tinggi".

#### 4.3. Transformation

Pada Transformation dilakukan perubahan data untuk pemodelan diantaranya membagi dataset untuk dataset latih dan uji.

Tabel 2. Transformasi

| Dataset      | Jumlah | Atribut |
|--------------|--------|---------|
| Dataset Awal | 443    | 9       |
| X_train      | 354    | 8       |
| Y_train      | 354    | 1       |
| X_test       | 89     | 8       |
| Y_test       | 89     | 1       |

Pembagian data dibagi kedalam 8 bagian untuk data training berjumlah 354 data dan 2 bagian untuk data testing atau berjumlah 89 data, Dengan proses pengacakan data sebanyak 42 kali untuk memvalidasi data dan meningkatkan nilai akurasi.

```
preprocessor = ColumnTransformer(
    transformers=[
        ('cat', OneHotEncoder(handle_unknown='ignore'),
         categorical_columns),
        ('num', StandardScaler(), numerical_columns)
    ])
```

Cuplikan kode di atas kolom string diubah dengan OneHotEncoder dan kolom numerik dengan StandardScaler. Keduanya digabung dalam parameter preprocessor menggunakan ColumnTransformer untuk transformasi berbeda pada DataFrame, untuk digunakan dalam pemodelan.

#### 4.4. Data Mining

Tahap ini Support Vector Machine akan diimplementasikan untuk mengklasifikasi menggunakan fungsi SVC (Support Vector Classifier). Model SVM didefinisikan dengan pipeline untuk menggabungkan langkah *preprocessing* dan klasifikasi menjadi satu alur kerja.

```
svm_model = Pipeline(steps=[
    ('preprocessor', preprocessor),
    ('classifier', SVC())
])
```

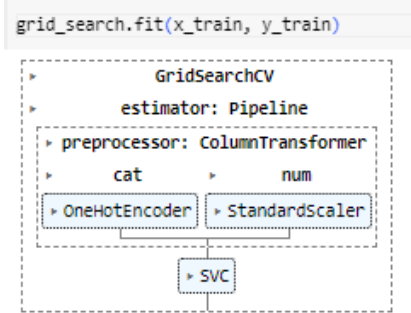
Selanjutnya, untuk Menentukan parameter dari SVC yang belum didefinisikan di atas digunakan fungsi GridSearchCV untuk menentukannya.

```
param_grid = {
    'classifier_kernel': ['linear', 'rbf'],
    'classifier_C': [0.1, 1, 10],
    'classifier_gamma': ['scale', 'auto'],
    'classifier_class_weight': [None, 'balanced']
}
skf = StratifiedKFold(n_splits=5, shuffle=True, random_state=42)
grid_search = GridSearchCV(svm_model, param_grid, cv=skf,
                           scoring='accuracy', n_jobs=-1)
```

Parameter grid yang dicoba meliputi berbagai nilai untuk parameter C, kernel, dan gamma (untuk kernel RBF), dan parameter yang ditemukan adalah C = 1, Class\_weight = balanced, Kernel = rbf dan gamma = scale.

Setelah menentukan parameter terbaik, model SVM dilatih menggunakan seluruh data yang telah diproses dengan parameter yang terpilih diastukan

menjadi satu alur kerja dengan pipeline dari mulai tahapan preprocessor hingga pemodelan.



Gambar 4. GridSearchCV dengan data latih

Setelah GridSearchCV selesai, objek `best_model` di bawah akan menyimpan model SVM yang telah dilatih dengan kombinasi parameter terbaik.

```
best_model = grid_search.best_estimator_
```

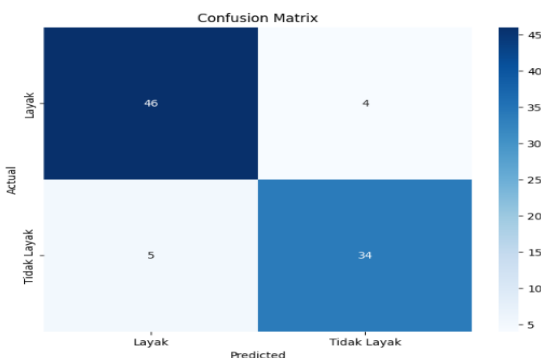
#### 4.5. Evaluation

Untuk mengukur kinerja model digunakan metrik evaluasi seperti akurasi, presisi, recall, dan F1-score memberikan wawasan mengenai seberapa baik model dapat mengklasifikasi penerima PKH berdasarkan data yang ada. Penting untuk dicatat bahwa evaluasi model dilakukan setelah data melalui tahap preprocessing, sehingga model dapat memberikan hasil yang lebih akurat dan relevan. Proses ini juga membantu dalam menilai kehandalan model dalam membuat klasifikasi.

```

plt.figure(figsize=(8, 6))
sns.heatmap(conf_matrix,
            annot=True, fmt='d',
            cmap='Blues',
            xticklabels=['Layak', 'Tidak Layak'],
            yticklabels=['Layak', 'Tidak Layak'])
plt.ylabel('Actual')
plt.xlabel('Predicted')
plt.title('Confusion Matrix')
plt.show()

```



Gambar 5. Confution Matrix

Keterangan :

1. True Negatives (TN): 34 (Model memprediksi 34 sampel negatif sebagai negatif)
2. True Positives (TP): 46 (Model memprediksi 46 sampel Positif yang sebenarnya positif).
3. False Positives (FP): 5 (Model memprediksi 5 sampel negatif sebagai positif)
4. False Negatives (FN): 4 (Model memprediksi 4 sampel positif sebagai negatif ).

Kemudian dilakukan perhitungan manual untuk hasil dari evaluasi confusion matrix

##### a. Accuracy

Accuracy adalah proporsi dari total prediksi yang benar :

$$\begin{aligned}
 \text{Akurasi} &= \frac{TP + TN}{TP + TN + FP + FN} \\
 &= \frac{46 + 34}{46 + 34 + 5 + 4} \\
 &= \frac{80}{89} = 0.898(89,89\%)
 \end{aligned}$$

##### b. Precision

Precision adalah proporsi prediksi positif yang benar :

$$\begin{aligned}
 \text{Precision} &= \frac{TP}{TP + FP} \\
 &= \frac{46}{46 + 5} = \frac{46}{51} = 0,901(90\%)
 \end{aligned}$$

##### c. Recall

Recall adalah proporsi dari total positif yang terdeteksi dengan benar

$$\begin{aligned}
 \text{Precision} &= \frac{TP}{TP + FN} \\
 \text{Recall} &= \frac{46}{46 + 4} \\
 &= \frac{46}{50} = 0.920(92\%)
 \end{aligned}$$

##### d. F1-Score

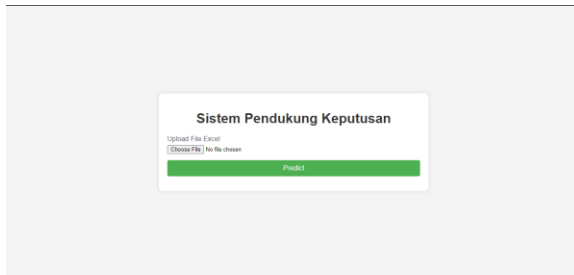
F1-Score adalah rata-rata dari precision dan recall.

$$\begin{aligned}
 F1 \text{ Score} &= 2x \frac{\text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}} \\
 &= 2x \frac{0.901 \times 0.920}{0.901 + 0.920} = 0.91(91\%)
 \end{aligned}$$

Hasil evaluasi model menunjukkan performa yang sangat baik dengan akurasi sebesar 89,8%, presisi yang sangat baik 90,%, recall 92%, dan fi-score 91%. Laporan klasifikasi dari confusion matrix menunjukkan bahwa model memiliki performa yang seimbang baik untuk kelas positif maupun negatif dengan nilai precision, recall dan f1-score.

#### 4.6. Deployment

Tahap ini memastikan bahwa model tidak hanya berfungsi dalam hal teknis namun juga dapat memberikan kepraktisan kepada pengguna.



Gambar 6. Tampilan web

### 5. KESIMPULAN DAN SARAN

Penggunaan Support Vector Machine dalam website untuk mengklasifikasi telah dilakukan dengan optimal. SVM mampu memberikan klasifikasi yang sangat baik mengenai status kelayakan penerima PKH, sehingga dapat digunakan untuk alat bantu dalam pengambilan keputusan. Dengan menerapkan SVM diperoleh nilai akurasi sebesar 89,89% menggunakan confusion matrix. Adapun untuk saran penelitian selanjutnya adalah tambahkan fitur atau faktor lain yang mungkin mempengaruhi terhadap kelayakan

#### DAFTAR PUSTAKA

- [1] R. Dewi and H. F. Andrianus, "Analisis pengaruh kebijakan bantuan langsung tunai (BLT) terhadap kemiskinan di indonesia periode 2005-2015," *MENARA Ilmu*, vol. 15, no. 2, pp. 77–84, 2021.
- [2] Badan Pusat Statistik Indonesia, "Profil Kemiskinan di Indonesia Maret 2023," *Badan Pusat statistik*, 2023. <https://www.bps.go.id/id/pressrelease/2023/07/17/2016/profil-kemiskinan-di-indonesia-maret-2023.html>
- [3] M. Jauharuddin, "Kecamatan Malo Berbasis Web Menggunakan Metode Naïve Bayes," vol. 2, pp. 73–77, 2022.
- [4] E. Fitriani, "Perbandingan Algoritma C4.5 Dan Naïve Bayes Untuk Menentukan Kelayakan Penerima Bantuan Program Keluarga Harapan," *Sistemasi*, vol. 9, no. 1, p. 103, 2020, doi: 10.32520/stmsi.v9i1.596.
- [5] I. P. Pratiwi, F. Ferdinandus, and A. D. Limantara, "Sistem Pendukung Keputusan Penerima Program Keluarga Harapan (PKH) Menggunakan Metode Simple Additive Weighting," vol. 8, no. 2, 2019.
- [6] A.I Pratiwi and Ahmad, "Implementasi Peraturan Menteri Sosial Nomor 1 Tahun 2018 Tentang Program Keluarga Harapan (PKH)," *Pendidik. Kim. PPs UNM*, vol. 2, no. 1, pp. 77–84, 2022.
- [7] A. Roihan, P. A. Sunarya, and A. S. Rafika, "Pemanfaatan Machine Learning dalam Berbagai Bidang: Review paper," *IJCIT (Indonesian J. Comput. Inf. Technol.)*, vol. 5, no. 1, pp. 75–82, 2020, doi: 10.31294/ijcit.v5i1.7951.
- [8] Suhardjono, W. Ganda, and H. Abdul, "Prediksi Kellusan Menggunakan Svm Berbasis Pso," *Bianglala Inform.*, vol. 7, no. 2, pp. 97–101, 2019.
- [9] N. Abizal, Maimun, and Yulindawati, "Efektivitas Program Keluarga Harapan (PKH) terhadap Kesejahteraan Masyarakat Masa Pandemi Covid-19 (Studi Kasus Kecamatan Tangan-tangan Kabupaten Aceh Barat Daya)," *J. Ilm. Basis Ekon. dan Bisnis*, vol. 1, no. 1, pp. 55–70, 2022, doi: 10.22373/jibes.v1i1.1576.
- [10] R. Nanda, E. Haerani, S. K. Gusti, and S. Ramadhani, "Klasifikasi Berita Menggunakan Metode Support Vector Machine," *J. Nas. Komputasi dan Teknol. Inf.*, vol. 5, no. 2, pp. 269–278, 2022, doi: 10.32672/jnkti.v5i2.4193.
- [11] F. Handayani *et al.*, "JEPIN (Jurnal Edukasi dan Penelitian Informatika) Komparasi Support Vector Machine, Logistic Regression Dan Artificial Neural Network dalam Prediksi Penyakit Jantung," *J. Edukasi dan Penelit. Inform.*, vol. 7, no. 3, p. Vol. 7 No. 3, 2021.
- [12] D. P. Utomo and M. Mesran, "Analisis Komparasi Metode Klasifikasi Data Mining dan Reduksi Atribut Pada Data Set Penyakit Jantung," *J. Media Inform. Budidarma*, vol. 4, no. 2, p. 437, 2020, doi: 10.30865/mib.v4i2.2080.
- [13] S. Anggraini, M. Akbar, A. Wijaya, H. Syaputra, and M. Sobri, "Klasifikasi Gejala Penyakit Coronavirus Disease 19 (COVID-19) Menggunakan Machine Learning," *J. Softw. Eng. Ampera*, vol. 2, no. 1, pp. 57–68, 2021, doi: 10.51519/journalsea.v2i1.105.
- [14] G. I. E. Soen, Marlina, and Renny, "Implementasi Cloud Computing dengan Google Colaboratory Pada Aplikasi Pengolah Data Zoom Participants," *J. Inform. Technol. Commun.*, vol. 6, no. 1, pp. 24–30, 2022.
- [15] Irmayanti, "Perancangan Sistem Informasi Penyewaan Thermoking Pada PT . Moderen Prima Transportasi Menggunakan Python Dengan Framework Flask," *JuSTICe*, vol. 1, no. 1, pp. 24–34, 2023.