

## Early Detection of GDM in First-Trimester Pregnant Women Using SVM Classifier

Achmad Baroqah Pohan<sup>1</sup>, Siti Masripah<sup>2</sup>, Lestari Yusuf<sup>3</sup>, Dini Nurlaela<sup>4</sup>, Lila Dini Utami<sup>5</sup>, Sri Wasiyanti<sup>6</sup>

<sup>1</sup> Information System, Sekolah Tinggi Teknologi Informasi NIIT, Jakarta, Indonesia

<sup>2</sup> Informatics, Universitas Bina Sarana Informatika, Bogor, Indonesia

<sup>3</sup> Informatics, Universitas Bina Sarana Informatika, Jakarta, Indonesia

<sup>4</sup> Information System, Universitas Bina Sarana Informatika, Jakarta, Indonesia

<sup>5</sup> Information System, Universitas Bina Sarana Informatika, Bogor, Indonesia

<sup>6</sup> Accounting Information System, Universitas Bina Sarana Informatika, Bogor, Indonesia

### ARTICLE INFORMATION

#### Artikel History::

Received: 18-05- 2026

Revised: 22-06- 2026

Accepted: 27-07- 2026

#### Keyword:

GDM  
Pregnant Women  
Classification  
SVM  
AUC

### ABSTRACT

Indonesia is among the top five countries in the world with the highest number of people living with diabetes, including those in the productive age group and pregnant women. Gestational Diabetes Mellitus (GDM) is a common metabolic complication during pregnancy and is generally diagnosed at 24–28 weeks of gestation. Delayed detection of GDM can limit opportunities for early intervention and increase the risk of complications for both the mother and the fetus. Therefore, an early detection approach is necessary, particularly during the first trimester of pregnancy. The development of Artificial Intelligence (AI), especially Machine Learning (ML), provides opportunities to develop disease prediction models based on clinical data. One algorithm that has demonstrated strong performance in classification tasks is the Support Vector Machine (SVM). This study aims to evaluate the performance of an SVM model using the Confusion Matrix and Area Under the Curve (AUC) in the early detection of GDM among first-trimester pregnant women. The research methodology includes the collection and preprocessing of clinical data, implementation of the algorithm, and testing or evaluation of model performance. The results show an accuracy of 71% and an AUC value of 76%. This study is expected to contribute to the early detection of GDM in pregnant women, thereby helping to reduce the risk of diabetes in children after birth.

### Corresponding Author:

Lestari Yusuf,  
Informatics,  
Universitas Bina Sarana Informatika,  
Jl. Kramat Raya No.98, RT.2/RW.9,  
Kwitang, Jakarta, Indonesia 10450  
Email: [lestari.lyf@bsi.ac.id](mailto:lestari.lyf@bsi.ac.id)

### INTRODUCTION

Indonesia is one of the 37 countries and territories in the Western Pacific Region. Indonesia has the fifth highest number of adults aged 20–79 years with diabetes in the world (International Diabetes Federational, 2025). Diabetes mellitus is not only experienced by individuals aged over 45 years, but is currently also found among children, young and productive age groups, including pregnant women (Fatrianti & Sulastri, 2025). In pregnant



women with diabetes mellitus, the body may experience difficulty adapting to hormonal changes during pregnancy, including insulin hormones. One characteristic of diabetes mellitus is an increase in blood sugar levels, and when the body cannot properly adapt to changes in insulin hormones (Septiana et al., 2024), it can lead to decreased insulin sensitivity, causing blood glucose levels to rise rapidly in some pregnant women and develop into Gestational Diabetes Mellitus (GDM) (Fatrianti & Sulastris, 2025).

Gestational Diabetes Mellitus (GDM) is one of the most common metabolic complications during pregnancy and is generally diagnosed through blood glucose examinations conducted between 24–28 weeks of gestation (Fatrianti & Sulastris, 2025) (Hu et al., 2023). Delayed examinations may result in a narrow time window for adequate intervention (Hu et al., 2023). Therefore, early detection of GDM in pregnant women is necessary and important.

Currently, many researchers have discussed GDM in pregnant women, both focusing specifically on the disease itself (Fatrianti & Sulastris, 2025) (Dewi et al., 2024) (Mato et al., 2023) (Rohmatulloh et al., 2024) and on GDM using advancements in Artificial Intelligence (AI), such as Machine Learning (ML) and Deep Learning (DL) (Bigdeli et al., 2025) (Septiana et al., 2024) (Hu et al., 2023). One of the algorithms commonly used is the Support Vector Machine (SVM) algorithm. SVM is capable of processing patients' clinical data and building predictive models to classify whether a person is likely to suffer from diabetes or not (Septiana et al., 2024). However, previous studies mainly discussed GDM in pregnant women during the second trimester, thus providing a significant opportunity to develop an intelligent model for predicting GDM in first-trimester pregnant women as an effort toward early prevention of Diabetes Mellitus (DM) progressing into Gestational Diabetes Mellitus (GDM).

Based on the background described above, the formulated research problems are: 1) How to perform early detection of GDM in pregnant women using Machine Learning, and 2) How to evaluate the results of the developed Machine Learning model. To address these problems, the objectives of this study are: 1) To use the Support Vector Machine (SVM) algorithm for early detection of GDM in pregnant women, and 2) To evaluate the developed model using a Confusion Matrix and Area Under the Curve (AUC). The expected benefits of this study are to contribute theoretically to the development of machine learning models and to assist medical personnel in the early detection of GDM in pregnant women.

## RESEARCH METHOD

This research was designed to address the identified problems through a systematic and structured scientific approach. The approach was carried out by implementing clear research stages, starting from problem formulation, data collection, analysis, and evaluation of the results, so that the findings obtained can be academically justified.

### 2.1. Data Source

The data used in this study were obtained from the Kaggle dataset website, namely [Kaggle Dataset: Early Prediction of GDM](#) in CSV format. The dataset taken from Kaggle is a Gestational Diabetes Mellitus (GDM) dataset of first-trimester pregnant women consisting of 10,000 records and 17 attributes, where one of the attributes is the diagnosis attribute with labels 1 (GDM) and 0 (No GDM).

### 2.2. Research Stages

Research stages are the levels or phases in a study that are carried out in a structured, sequential, standardized, logical, and systematic manner (Azis, 2023).



Figure 2. Research Stages

This research design represents the initial framework before entering the implementation stage of the Support Vector Machine (SVM) method. It explains the process of classifying early gestational diabetes disease using the SVM method. The first stage is data collection, followed by the data processing stage, where the data are divided into two categories: training data and testing data. After the data division process, the study proceeds to the SVM modeling stage. Subsequently, the model evaluation process is carried out using a confusion matrix, precision calculation, F1-score, and ROC analysis. The final stage consists of presenting the results, conclusions, and discussion of the research. The purpose of this study is to perform early detection of whether a pregnant woman is affected by gestational diabetes or not.

### 2.2.1 Data Collection

The data were obtained from secondary or public data available on the Kaggle website, consisting of approximately 6,500 positive GDM cases (65%) and 3,500 negative GDM cases (35%). This dataset simulates realistic patient profiles based on established epidemiological relationships and clinical guidelines for conducting GDM screening. The dataset attributes are presented in Table 1:

Table 1. GDM Attributes

| No | Attribute                      | Description                                                         | Range and Values                           |
|----|--------------------------------|---------------------------------------------------------------------|--------------------------------------------|
| 1  | <i>patient_id</i>              | Identitas Unik Pasien                                               | P0001 - P10000                             |
| 2  | <i>age</i>                     | Usia Ibu Hamil Pada Saat Kunjungan Pendaftaran                      | 18 - 45 years                              |
| 3  | <i>pre_pregnancy_bmi</i>       | Indeks Massa Tubuh Sebelum Kehamilan (Kg/m2)                        | 16 - 45 kg/m <sup>2</sup>                  |
| 4  | <i>ethnicity</i>               | Latar Belakang Etnis yang dilaporkan                                | Caucasian, Asian, Hispanic, African, Other |
| 5  | <i>family_history_dm</i>       | Keluarga tingkat pertama yang menderita Diabetes Tipe 1 atau Tipe 2 | 0 = No, 1 = Yes                            |
| 6  | <i>previous_gdm</i>            | Riwayat GDM Pada Kehamilan Sebelumnya                               | 0 = No, 1 = Yes                            |
| 7  | <i>pcos</i>                    | Didiagnosis menderita sindrom ovarium polikistik                    | 0 = No, 1 = Yes                            |
| 8  | <i>previous_macrosomia</i>     | Riwayat Persalinan bayi dengan berat diatas 4 KG                    | 0 = No, 1 = Yes                            |
| 9  | <i>booking_gestational_age</i> | Usia Kehamilan pada Kunjungan Prenatal pertama                      | 6 - 16 weeks                               |
| 10 | <i>early_rbs_mgdl</i>          | Kadar Glukosa darah tanpa puasa pada kunjungan trimester pertama    | 70 - 200 mg/dL                             |
| 11 | <i>early_ppbs_mgdl</i>         | Kadar Glukosa darah 2 jam setelah makan                             | 90 - 250 mg/dL                             |
| 12 | <i>early_hb1c_percent</i>      | Persentase Hemoglobin Terглиkasi                                    | 4.0 - 8.0 %                                |
| 13 | <i>early_ogtt_performed</i>    | Apakah Tes OGTT 75g telah dilakukan pada trimester pertama ?        | 0 = No, 1 = Yes                            |
| 14 | <i>early_ogtt_fasting_mgdl</i> | Kadar Glukosa Plasma Puasa selama OGTT 75g                          | 70 - 140 mg/dL                             |
| 15 | <i>early_ogtt_1h_mgdl</i>      | Kadar Glukosa plasma 1 jam setelah pemberian beban glukosa 75g      | 100 - 250 mg/dL                            |
| 16 | <i>early_ogtt_2h_mgdl</i>      | Kadar Glukosa plasma 2 jam setelah pemberian beban glukosa 75g      | 90 - 220 mg/dL                             |
| 17 | <i>gdm_outcome</i>             | Diagnosa Diabetes Melitus Gestasional telah dikonfirmasi            | 0 = No GDM, 1 = GDM                        |

### 2.2.2 Data Pre-Processing

The collected data were subsequently processed to produce the best possible data, thereby providing sufficiently accurate results (Prasetyo et al., 2024). Several processing stages were carried out, including outlier cleaning, missing value imputation, data adjustment and labeling, data normalization, as well as the division of training and testing data. The following are the stages in data preprocessing:

#### 1. Outlier Cleaning

Outliers are extremely unusual or abnormal values within a dataset. Some of the benefits of outlier cleaning include reducing the influence of outliers on the model, improving robustness, increasing accuracy, and enhancing stability during training. The implementation of outlier cleaning for each parameter used can be seen in Figure 1 below.

| patient_id | age  | pre_pregnancy_bmi | ethnicity | family_history_dm | previous_gdm | pcos | previous_macrosomia | booking_gestational_age | early_rbs_mgdl | early_ppbs_mgdl | early |
|------------|------|-------------------|-----------|-------------------|--------------|------|---------------------|-------------------------|----------------|-----------------|-------|
| 804        | 30.0 | 31.4              | Asian     | 0                 | 0            | 0    | 0                   | 3.4                     | 122.7          | 179.7           | 0     |
| 801        | 30.0 | 30.0              | Caucasian | 0                 | 0            | 1    | 1                   | 3.1                     | 119.0          | 168.4           | 0     |
| 807        | 30.0 | 31.2              | African   | 0                 | 0            | 0    | 0                   | 3.2                     | 124.5          | 168.0           | 0     |
| 1443       | 31.0 | 30.0              | Asian     | 1                 | 0            | 0    | 0                   | 3.0                     | 128.8          | 173.0           | 0     |
| 9187       | 31.0 | 31.2              | Asian     | 1                 | 0            | 0    | 0                   | 3.0                     | 121.0          | 173.0           | 0     |
| 3286       | 32.0 | 31.2              | Caucasian | 0                 | 0            | 0    | 0                   | 7.2                     | 128.7          | 168.1           | 0     |
| 3614       | 32.0 | 31.2              | Hispanic  | 1                 | 0            | 0    | 0                   | 15.0                    | 128.7          | 168.1           | 0     |
| 4054       | 32.0 | 25.5              | African   | 0                 | 0            | 0    | 0                   | 10.8                    | 112.8          | 160.5           | 0     |
| 3178       | 32.0 | 38.1              | Caucasian | 0                 | 0            | 0    | 0                   | 3.6                     | 122.3          | 167.7           | 0     |
| 9082       | 32.0 | 18.9              | Caucasian | 0                 | 0            | 0    | 0                   | 7.2                     | 98.8           | 128.0           | 0     |
| 4081       | 32.0 | 34.4              | Asian     | 0                 | 0            | 1    | 0                   | 3.9                     | 110.0          | 173.9           | 0     |
| 8887       | 32.0 | 36.0              | Other     | 0                 | 0            | 0    | 0                   | 11.0                    | 118.0          | 162.7           | 0     |
| 7878       | 33.0 | 23.1              | Caucasian | 0                 | 0            | 0    | 0                   | 11.0                    | 103.8          | 160.0           | 0     |
| 7073       | 33.0 | 33.8              | Caucasian | 0                 | 0            | 0    | 0                   | 3.9                     | 110.8          | 160.0           | 0     |
| 8073       | 33.0 | 33.0              | Asian     | 0                 | 0            | 0    | 0                   | 3.1                     | 88.0           | 168.1           | 0     |
| 8789       | 33.0 | 32.1              | Asian     | 1                 | 0            | 0    | 0                   | 8.4                     | 120.4          | 168.4           | 0     |
| 9183       | 34.0 | 31.4              | Asian     | 0                 | 0            | 0    | 0                   | 3.2                     | 140.0          | 174.0           | 0     |
| 9487       | 34.0 | 38.0              | Caucasian | 0                 | 0            | 0    | 1                   | 11.7                    | 144.0          | 210.0           | 0     |

Figure 2. Outlier Cleaning on the Data

## 2. Missing Data Value Imputation

Missing data value imputation was carried out by filling in missing data using the mean and median values. The determination of whether to use mean or median imputation considered several factors. Columns filled with mean values generally have a normal data distribution and relatively few missing values. Meanwhile, columns filled with median values usually have an abnormal data distribution and may contain a large number of missing values. Missing data value imputation provides several benefits, including maintaining sample size, improving data usability, and reducing bias (Syahputra et al., 2023). The following figure illustrates the implementation of missing data value imputation:

```

df['early_ogtt_fasting_mgdl'] = df['early_ogtt_fasting_mgdl'].replace(0,np.median(df['early_ogtt_fasting_mgdl']))
df['early_ogtt_1h_mgdl'] = df['early_ogtt_1h_mgdl'].replace(0,np.median(df['early_ogtt_1h_mgdl']))
df['early_ogtt_2h_mgdl'] = df['early_ogtt_2h_mgdl'].replace(0,np.median(df['early_ogtt_2h_mgdl']))

df.head()

```

| Index | patient_id | age  | pre_pregnancy_bmi | ethnicity | family_history_dm | previous_gdm | pcos | previous_macrosomia | booking_gestational_age | early_rbs_mgdl | early_ppbs_mgdl | early_hba1c_percent | early_ogtt_performed | early |
|-------|------------|------|-------------------|-----------|-------------------|--------------|------|---------------------|-------------------------|----------------|-----------------|---------------------|----------------------|-------|
| 0     | P0001      | 30.5 | 25.4              | Caucasian | 0                 | 0            | 0    | 0                   | 10.6                    | 117.5          | 122.5           | 5.26                | 0                    |       |
| 1     | P0002      | 32.8 | 24.3              | Asian     | 0                 | 0            | 0    | 0                   | 9.5                     | 101.0          | 144.0           | 5.27                | 0                    |       |
| 2     | P0003      | 27.9 | 18.8              | Caucasian | 0                 | 1            | 0    | 0                   | 15.6                    | 127.4          | 147.5           | 6.1                 | 0                    |       |
| 3     | P0004      | 25.3 | 19.5              | Caucasian | 1                 | 0            | 1    | 0                   | 16.3                    | 135.5          | 163.7           | 5.47                | 0                    |       |
| 4     | P0005      | 22.5 | 18.2              | Caucasian | 0                 | 0            | 1    | 0                   | 11.8                    | 108.4          | 106.7           | 5.38                | 0                    |       |

Figure 3. Missing Data Value Imputation

## 3. Data Normalization

Data normalization is a process in which the values within a dataset are adjusted or normalized to have a uniform or comparable scale (Septiana et al., 2024). The data columns to which normalization was applied are as follows: *age*, *pre\_pregnancy\_bmi*, *ethnicity*, *family\_history\_dm*, *previous\_gdm*, *pcos*, *previous\_macrosomia*, *booking\_gestational\_age*, *early\_rbs\_mgdl*, *early\_ppbs\_mgdl*, *early\_hba1c\_percent*, *early\_ogtt\_performed*. However, for the three columns, namely: *early\_ogtt\_fasting\_mgdl*, *early\_ogtt\_1h\_mgdl*, *early\_ogtt\_2h\_mgdl*, because the amount of missing data exceeded 50% of the total dataset records.

```

features = [
    'age',
    'pre_pregnancy_bmi',
    'ethnicity',
    'family_history_dm',
    'previous_gdm',
    'pcos',
    'previous_macrosomia',
    'booking_gestational_age',
    'early_rbs_mgdl', 'early_ppbs_mgdl', 'early_hba1c_percent', 'early_ogtt_performed'
]

X = df[features].copy()
y = df['gdm_outcome']

```

Figure 4. Features Used as Columns and Labels

Data normalization provides several advantages, including improving model convergence, reducing the influence of outliers, adjusting data to the same scale, and enhancing model performance and accuracy. The implementation of data normalization is shown in the following figure:



Figure 5. Implementation of Data Normalization

#### 4. Training and Testing Data Split

This division aims to evaluate how well the model trained on the training set can generalize to previously unseen data contained in the testing set. The implementation is shown in Figure 6 below:



Figure 6. Training Data and Testing Data Split

#### 2.2.3. Analysis and Implementation of the Machine Learning Model

After the data were processed, further analysis was conducted using Cross Validation to apply the algorithm used in this context, namely Support Vector Machine (SVM). SVM is one of the most appropriate discriminative methods used for classification (Karinari & Badriyah, 2020). This method works based on structural risk minimization, which aims to find the best hyperplane that separates two classes in the input space. The SVM kernel used in this study is the Gaussian RBF kernel, as shown in Equation 1 (Jordan et al., 2026):

$$K \left( \vec{x}_i, \vec{x}_j \right) = \exp \left[ -\frac{\left\| \vec{x}_i - \vec{x}_j \right\|^2}{2\sigma^2} \right] \quad \text{Equation 1}$$

#### 2.2.4 Evaluation

The evaluation phase focuses on measuring the performance of a prediction or classification model using relevant metrics (Arfian et al., 2025). The SVM model will be evaluated using a confusion matrix, learning curve, and classification report, which are divided into several components, namely Accuracy, Precision, Recall, F1-

Score, and Support. A confusion matrix generally has the form of a square table with four main cells:

1. **True Positive (TP)**, namely the number of data instances correctly predicted as positive by the model.
2. **True Negative (TN)**, namely the number of data instances correctly predicted as negative by the model.
3. **False Positive (FP)**, namely the number of data instances that are actually negative but incorrectly predicted as positive by the model.
4. **False Negative (FN)**, namely the number of data instances that are actually positive but incorrectly predicted as negative by the model.

The confusion matrix provides more detailed information about the performance of a classification model, enabling the calculation of evaluation metrics such as accuracy, precision, recall, and others. The following are several evaluation metrics that can be calculated using a confusion matrix (James et al., 2017):

1. **Accuracy**

Accuracy measures the extent to which the model can make correct predictions overall. The formula is:

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN}$$

2. **Precision**

Precision measures the extent to which the positive predictions made by the model are correct. The formula is:

$$Precision = \frac{TP}{TP + FP}$$

3. **Recall (Sensitivity or True Positive Rate)**

Recall measures the extent to which the model can correctly identify the actual positive class. The formula is:

$$Recall = \frac{TP}{TP + FN}$$

4. **F1-Score**

The F1-score is a composite measure that combines precision and recall. It is useful when seeking a balance between precision and recall. The formula is:

$$F1-Score = 2 \times \frac{(Precision \times Recall)}{(Precision + Recall)}$$

The accuracy obtained from the model will serve as a reference for future studies to conduct comparisons in order to determine the best-performing model. The best model will then proceed to the deployment stage.

## RESULTS AND DISCUSSION

The implementation of SVM along with its best parameters is explained in Figure 7 below.

```

svm_pipeline = Pipeline([
    ('preprocessing', preprocessor),
    ('svm', SVM(
        kernel='rbf',
        C=1,
        gamma='scale',
        probability=True,
        class_weight='balanced',
        random_state=42
    ))
])

svm_pipeline.fit(x_train, y_train)

```

The Pipeline diagram shows a 'preprocessing: ColumnTransformer' step with two parallel paths: 'num' containing 'StandardScaler' and 'cat' containing 'OneHotEncoder'. Both paths lead to an 'svm' step.

Figure 7. Implementation of the SVM Algorithm Model with Its Parameters

Figure 7 illustrates the pipeline script (step-by-step process) of the SVM algorithm. StandardScaler is used for numerical data normalization by standardizing all data values. OneHotEncoder is used to convert categorical data, in this case the *Ethnic* attribute, into numerical type data.

In the SVM configuration, the kernel used is the RBF (Radial Basis Function) type to capture possible non-linear relationships among variables. The parameters applied in the script are as follows:

1. **C = 1**, which regulates the balance between maximum margin and classification error penalty.
2. **gamma = 'scale'**, which automatically adjusts the gamma value based on the data distribution.
3. **class\_weight = 'balanced'**, to address class distribution imbalance by assigning proportional weights according to the frequency of each class.
4. **probability = True**, so that the model can generate probability estimates required for ROC and AUC-based evaluation.

Next, the confusion matrix of the SVM model is calculated. The results are explained in Figure 8.

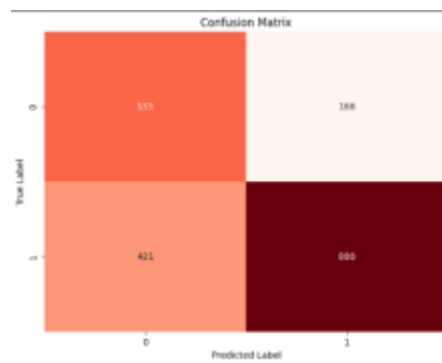


Figure 8. Confusion Matrix

Figure 8 explains the confusion matrix results with the values of True Negative (TN) = 533, False Positive (FP) = 166, False Negative (FN) = 421, and True Positive (TP) = 880. This confusion matrix shows the distribution of the model's predictions against the actual classes in the testing data. Overall, the model successfully classified 533 non-GDM samples correctly (True Negative). However, there were 166 cases that were actually negative but predicted as positive (False Positive). This indicates that a small proportion of individuals without GDM were identified as being at risk by the model.

The specificity (true negative rate) value can be calculated using Equation 5 as follows:

$$Specificity = \frac{TN}{TN+FP} = \frac{533}{533+166} \approx 0.76 \quad \text{Equation 5}$$

This calculation means that the model is capable of correctly identifying approximately 76% of negative cases.

```
from sklearn.metrics import classification_report
print(classification_report(y_test, y_pred))
```

|              | precision | recall | f1-score | support |
|--------------|-----------|--------|----------|---------|
| 0            | 0.56      | 0.76   | 0.64     | 699     |
| 1            | 0.84      | 0.68   | 0.75     | 1301    |
| accuracy     |           |        | 0.71     | 2000    |
| macro avg    | 0.70      | 0.72   | 0.70     | 2000    |
| weighted avg | 0.74      | 0.71   | 0.71     | 2000    |

Figure 9. Classification Report

Next, the ROC AUC graph is presented to determine the extent of the model's ability to distinguish between the two classes. The ROC and AUC results are explained in Figure 10.

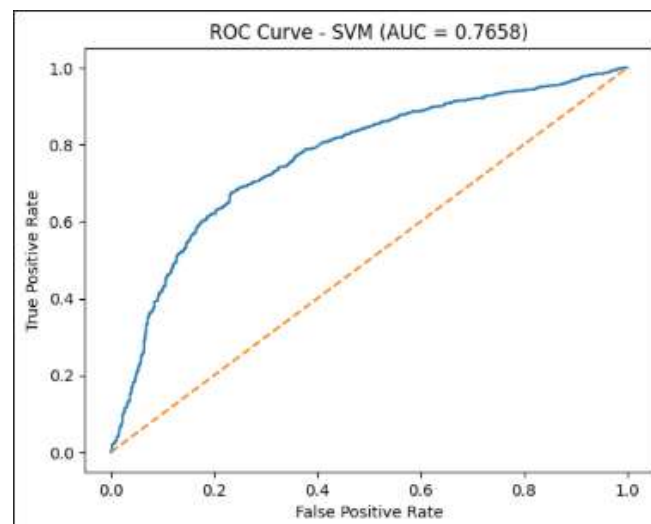


Figure 10. ROC Curve Graph

Figure 10 displays the ROC (Receiver Operating Characteristic) curve for the Support Vector Machine (SVM) model with an AUC (Area Under the Curve) value of 0.7658. The explanation is as follows:

1. The ROC curve illustrates the relationship between: True Positive Rate (TPR) or Sensitivity/Recall on the Y-axis, and False Positive Rate (FPR) on the X-axis for various classification threshold values.
2. The dashed diagonal line represents the performance of a random classifier with an AUC value of 0.5. A good model should have a curve positioned above this diagonal line.
3. The graph shows that the ROC curve of the SVM model consistently lies above the diagonal line, indicating that the model has better discriminative ability compared to random classification.
4. The AUC value of 0.7658 indicates that the model has a good classification capability (good discrimination). Probabilistically, this value can be interpreted as approximately a 76.58% chance that the model will assign a higher score to a positive sample than to a randomly selected negative sample.

The curve shape, which bends relatively toward the upper-left corner, indicates that the model is able to achieve a sufficiently high sensitivity level while maintaining a controlled increase in FPR, although it has not yet reached optimal performance (AUC approaching 1).

## CONCLUSION

The evaluation results indicate that the model has adequate discriminative capability in identifying the risk of GDM, making it potentially useful as a decision support system for the early screening of pregnant women.

Nevertheless, there is still room for improvement, particularly in enhancing the model's sensitivity to minimize classification errors in positive cases.

For future research, it is recommended to conduct comparisons with other machine learning algorithms, perform more in-depth hyperparameter optimization, and consider the use of feature selection or ensemble learning techniques to improve predictive performance. In addition, validation using real clinical datasets from healthcare institutions is important to ensure the generalizability and practical implementation of the model in real-world contexts.

## REFERENCES

- Arfian, A., Siregar, J., & Wijayanti, D. (2025). INTEGRASI MODEL NEURAL NETWORK DENGAN SVM DAN KNN UNTUK DETEKSI DINI PENYAKITA JANTUNG PADA PENDERITA HYPERTENSI. *Technologia : Jurnal Ilmiah*, 16(1), 104. <https://doi.org/10.31602/tji.v16i1.17046>
- Azis, Y. A. (2023). 7 Tahapan Penelitian yang Wajib Kamu Ketahui. [https://deepublishstore.com/blog/penelitian-skripsi/tahapan-penelitian/#:~:text=Tahapan penelitian adalah level atau,baku%2C logis dan juga sistematis](https://deepublishstore.com/blog/penelitian-skripsi/tahapan-penelitian/#:~:text=Tahapan%20logis%20dan%20juga%20sistematis).
- Bigdeli, S. K., Ghazisaedi, M., Ayyoubzadeh, S. M., Hantoushzadeh, S., & Ahmadi, M. (2025). Predicting Gestational Diabetes Mellitus in the first trimester using machine learning algorithms: a cross-sectional study at a hospital fertility health center in Iran. *BMC Medical Informatics and Decision Making*, 25(1), 3. <https://doi.org/10.1186/s12911-024-02799-3>
- Dewi, R. S., Martini, S., & Isfandiari, M. A. (2024). Prevalence and risk factors for gestational diabetes in Jambi, Indonesia. *African Journal of Reproductive Health*, 28(10 Special Edition), 118–124. <https://doi.org/10.29063/ajrh2024/v28i10s.14>
- Fatrianti, I., & Sulastri. (2025). Analisis Predisposisi Kejadian Diabetes Mellitus Gestasional pada Ibu Hamil Trisemester 3 di Klinik Pratama Dokter Abdul Radjak Cengkareng. *MAHESA : Malahyati Health Student Journal*, 5(8), 3461–3470. <https://doi.org/10.33024/>
- Hu, X., Hu, X., Yu, Y., & Wang, J. (2023). Prediction model for gestational diabetes mellitus using the XG Boost machine learning algorithm. *Frontiers in Endocrinology*, 14(March), 1–10. <https://doi.org/10.3389/fendo.2023.1105062>
- International Diabetes Federational. (2025). *Indonesia - Laporan negara diabetes 2000 — 2050*. Diabetesatlas.Org. <https://diabetesatlas.org/data-by-location/country/indonesia/>
- James, G., Witten, D., Hastie, T., & Tibshirani, R. (2017). *An introduction to statistical learning: with applications in R*. Springer : Springer Science+Business Media. [https://www.stat.berkeley.edu/~rabbee/s154/ISLR\\_First\\_Printing.pdf](https://www.stat.berkeley.edu/~rabbee/s154/ISLR_First_Printing.pdf)
- Jordan, M., Kleinberg, J., & Schölkopf, B. (2026). *Pattern Recognition and Machine Learning*. <https://www.microsoft.com/en-us/research/wp-content/uploads/2006/01/Bishop-Pattern-Recognition-and-Machine-Learning-2006.pdf>
- Karinasari, I., & Badriyah, T. (2020). DETEKSI DINI PENYAKIT IUGR (INTRA UTERINE GROWTH RETRICTION) DENGAN METODE SVM (SUPPORT VECTOR MACHINE). *Kumpulan Jurnal Ilmu Komputer (KLIK)*, 07.
- Mato, R., Aspilayuli, & Suhartatik. (2023). Literature Review : Faktor Yang Mempengaruhi Diabetes Mellitus Gestasional. *JIMPK : Jurnal Ilmiah Mahasiswa & Penelitian Keperawatan*, 3(4), 111–120.
- Prasetyo, B. R., Wahyuni, E. D., & Kusumantara, P. M. (2024). KOMPARASI PERFORMA MODEL BERBASIS ALGORITMA RANDOM FOREST DAN LIGHTGBM DALAM MELAKUKAN KLASIFIKASI DIABETES MELITUS GESTASIONAL. *Jurnal Informatika Dan Teknik Elektro Terapan*, 12(3). <https://doi.org/10.23960/jitet.v12i3.4817>
- Rohmatulloh, V. R., Riskiyah, Pardjianto, B., & Kinasih, L. S. (2024). Hubungan Usia dan Jenis Kelamin Terhadap Angka Kejadian Diabetes Melitus Tipe 2 Berdasarkan 4 Kriteria Diagnosis di Poliklinik Penyakit Dalam RSUD Karsa Husada Kota Batu. *PREPOTIF: Jurnal Kesehatan Masyarakat*, 8(1), 2528–2543. <https://journal.universitaspahlawan.ac.id/index.php/prepotif/article/view/27198/19343>
- Septiana, T., Muda, M. A., Budiyanto, D., Pratama, M., & Jaya, W. P. (2024). Analisis Penggunaan Support Vector Machine pada Deteksi Dini Penyakit Diabetes Melitus. *Jurnal Penelitian Inovatif*, 4(3), 1631–1640. <https://doi.org/10.54082/jupin.643>
- Syahputra, H., Naibaho, S. I., Maulana, M. A., Zulfahmi, I., & Sinaga, E. P. (2023). Perbandingan Algoritma Support Vector Machine (SVM) dan Decision Tree Untuk Deteksi Tingkat Depresi Mahasiswa. *BINA INSANI ICT JOURNAL*, 10(1), 52–61.