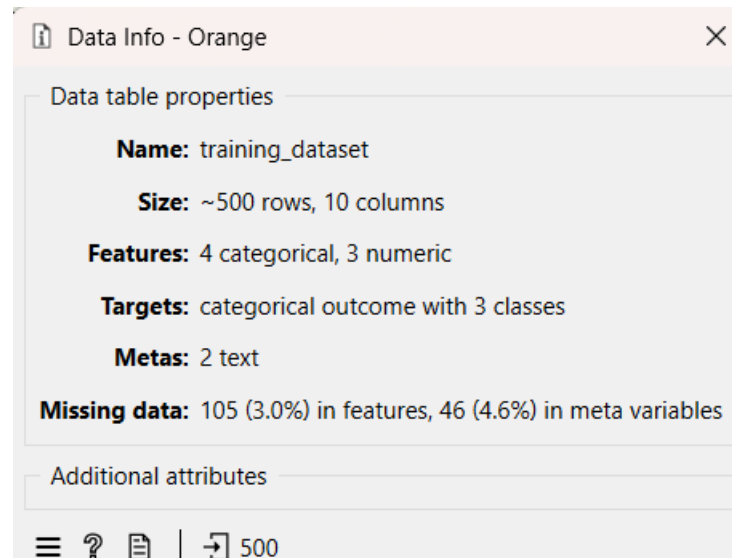


Soal 1 klasifikasi

A. Data Understanding

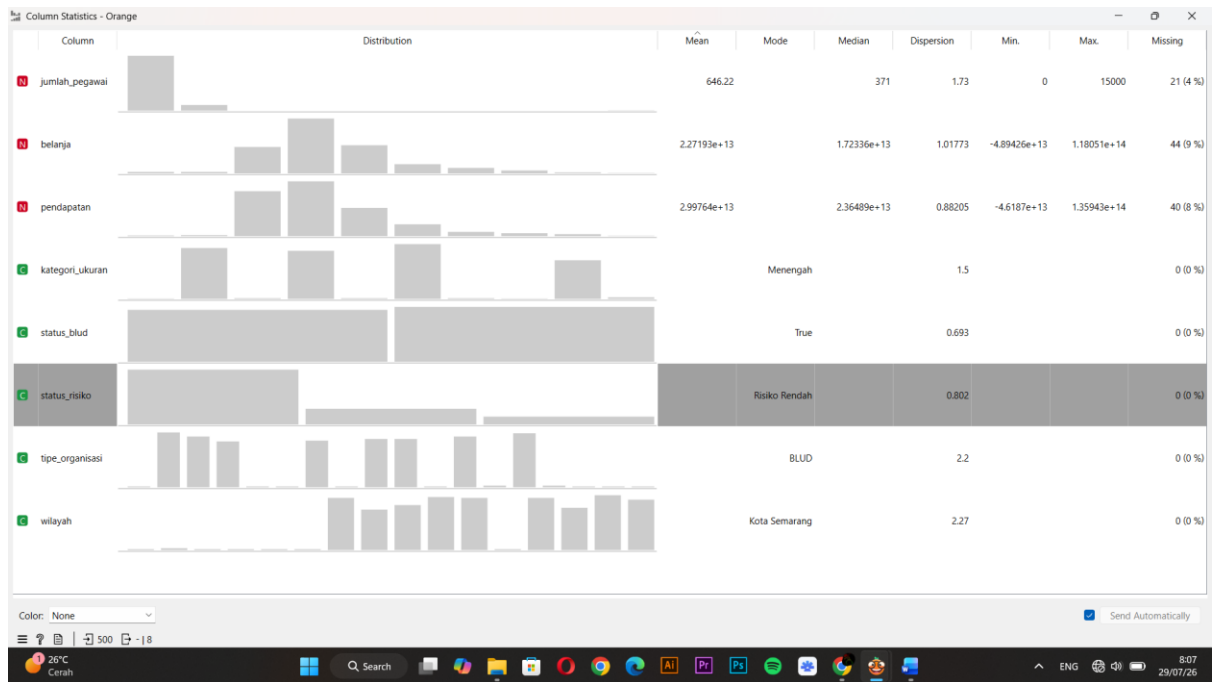
Identifikasi



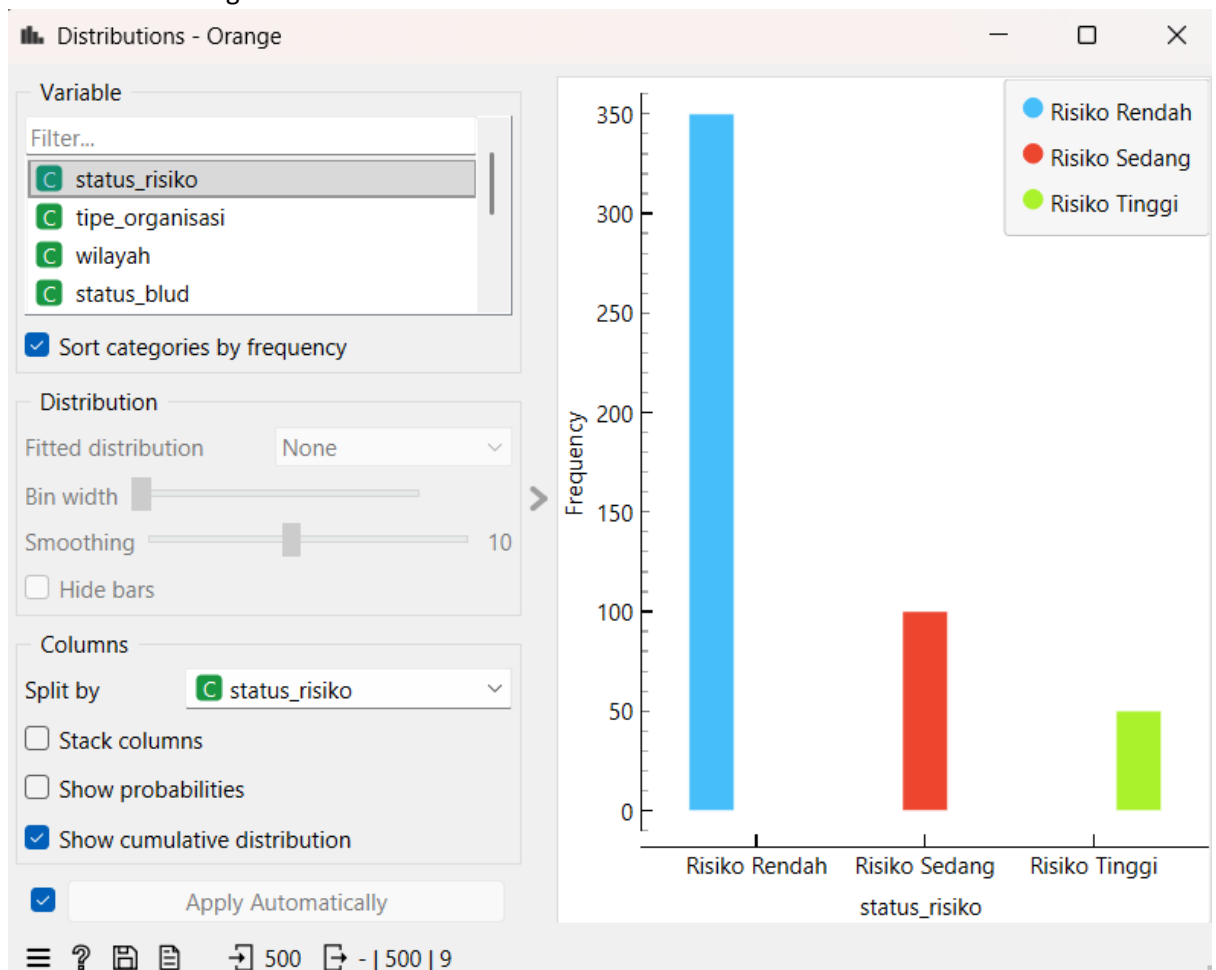
- a) Jumlah record = 500
- b) Jumlah atribut/variabel/feature = 10
- c) Tipe masing-masing atribut, d) atribut feature, e) atribut target (class) =

Columns (Double click to edit)			
	Name	Type	Role
1	tipe_organisasi	C categorical	feature
2	wilayah	C categorical	feature
3	status_blud	C categorical	feature
4	kategori_ukuran	C categorical	feature
5	jumlah_pegawai	N numeric	feature
6	pendapatan	N numeric	feature
7	belanja	N numeric	feature
8	status_risiko	C categorical	target
9	kode_organisasi	S text	meta
10	tanggal_pembe...	S text	meta

Statistik deskriptif



Distribusi kelas target

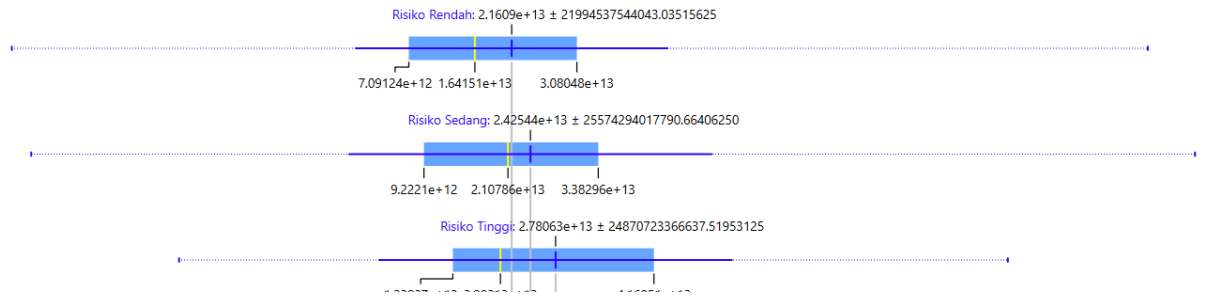


Identifikasi

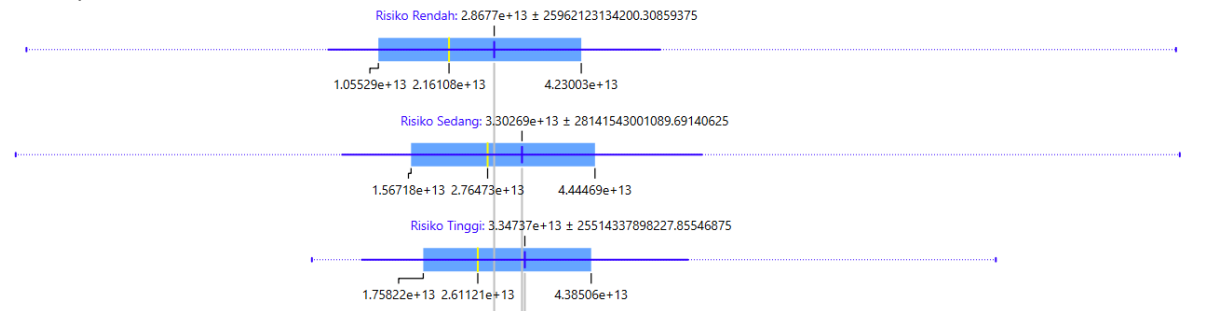
a) Missing value = ada, 105 in features dan 46 in meta

- b) Class imbalance = ada, risiko rendah paling banyak, jauh dibanding risiko sedang dan risiko tinggi
- c) Outlier = ada, terlihat di box plot ada titik yang sangat jauh dari sebaran normal jumlah pegawai untuk risiko rendah, begitu juga dengan feature numerik lain, terdapat titik yang melebar dari sebaran normal

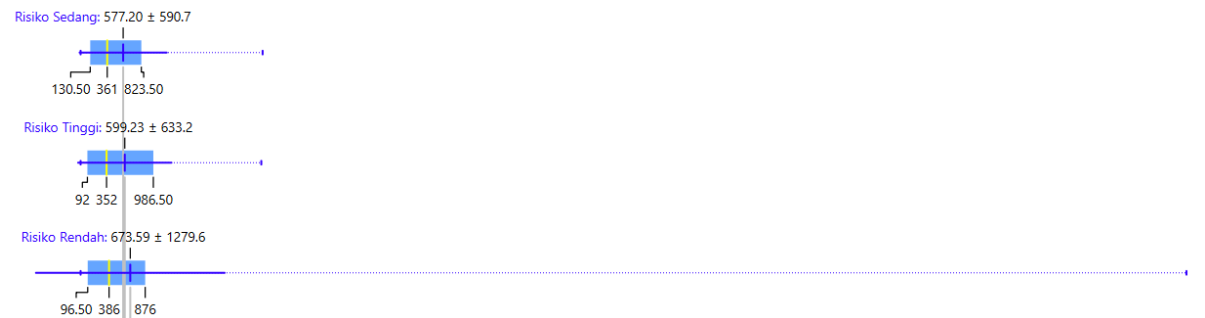
belanja



Pendapatan



Jumlah_pegawai



- d) Atribut tidak relevan = ada, terlihat di rank atribut status_blud memiliki impact terhadap target relatif terlalu sedikit.

Rank - Orange

Scoring Methods

- ☒ Information Gain
- ☒ Information Gain Ratio
- ☒ Gini Decrease
- ☐ ANOVA
- ☐ χ^2
- ☐ ReliefF
- ☐ FCBF

Select Attributes

- ☐ None
- ☐ All
- ☐ Manual
- ☒ Best ranked: 5

☒ Send Automatically

		#	Info. gain	Gain ratio	Gini
1	C tipe_organisasi	18	0.059	0.019	0.023
2	C wilayah	16	0.038	0.012	0.012
3	N pendapatan		0.017	0.008	0.007
4	C kategori_ukuran	10	0.014	0.007	0.004
5	N belanja		0.011	0.006	0.004
6	N jumlah_pegawai		0.009	0.004	0.003
7	C status_blud	2	0.001	0.001	0.000

Missing values will be imputed as needed.

Identifikasi atribut yang relevan

Berdasarkan rank, atribut yang relevan adalah tipe_organisasi, wilayah, pendapatan, kategori_ukuran, belanja dan jumlah_pegawai. Status_blud memiliki impact terhadap target relatif terlalu sedikit.

B. Data preparation

1. Menangani missing value

Impute - Orange

Default Method

☐ Don't impute
 ☐ Model-based imputer (simple tree)

☒ Average/Most frequent
 ☐ Random values

☐ As a distinct value
 ☐ Remove instances with unknown values

☐ Fixed values; numeric variables: , time:

Individual Attribute Settings

Filter...

C	tipe_organisasi -> average
C	wilayah -> average
C	status_blud -> average
C	kategori_ukuran -> average
N	jumlah_pegawai -> average
N	pendapatan -> average
N	belanja -> average
C	status_risiko -> average

☐ Default (above)
 ☐ Don't impute
 ☒ Average/Most frequent
 ☐ As a distinct value
 ☐ Model-based imputer (simple tree)
 ☐ Random values
 ☐ Remove instances with unknown values
 ☐ Fixed value

☒

500 | - 500

2. Menangani outlier

Outliers - Orange

Method

Parameters

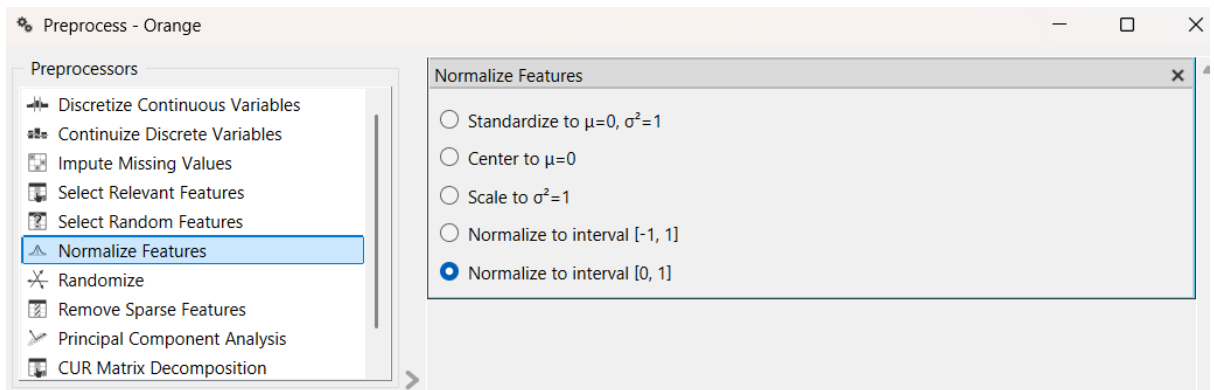
Contamination:

☐ Replicable training

☒

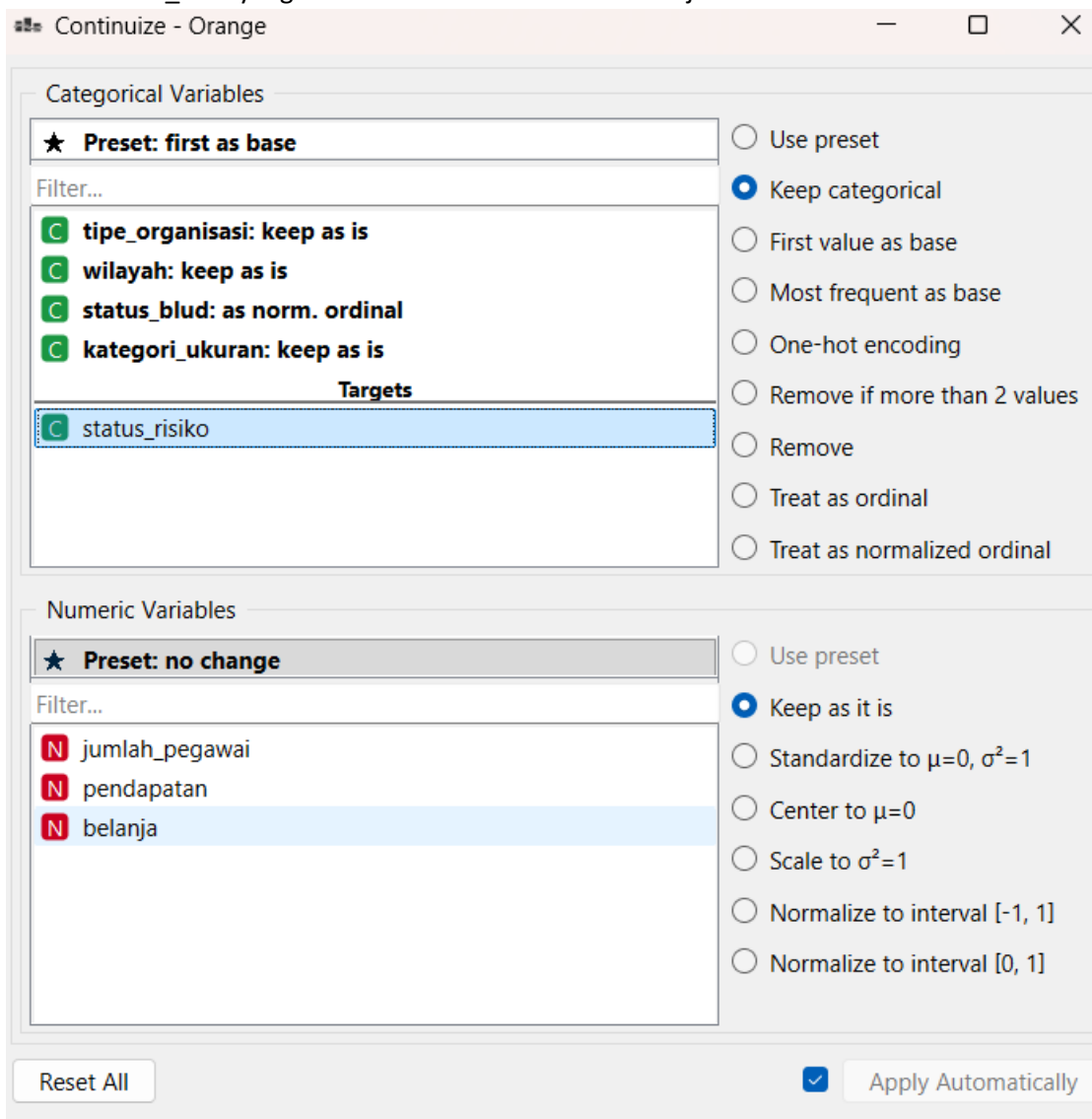
500 | 450 | 50 | 500

3. Normalisasi/scaling

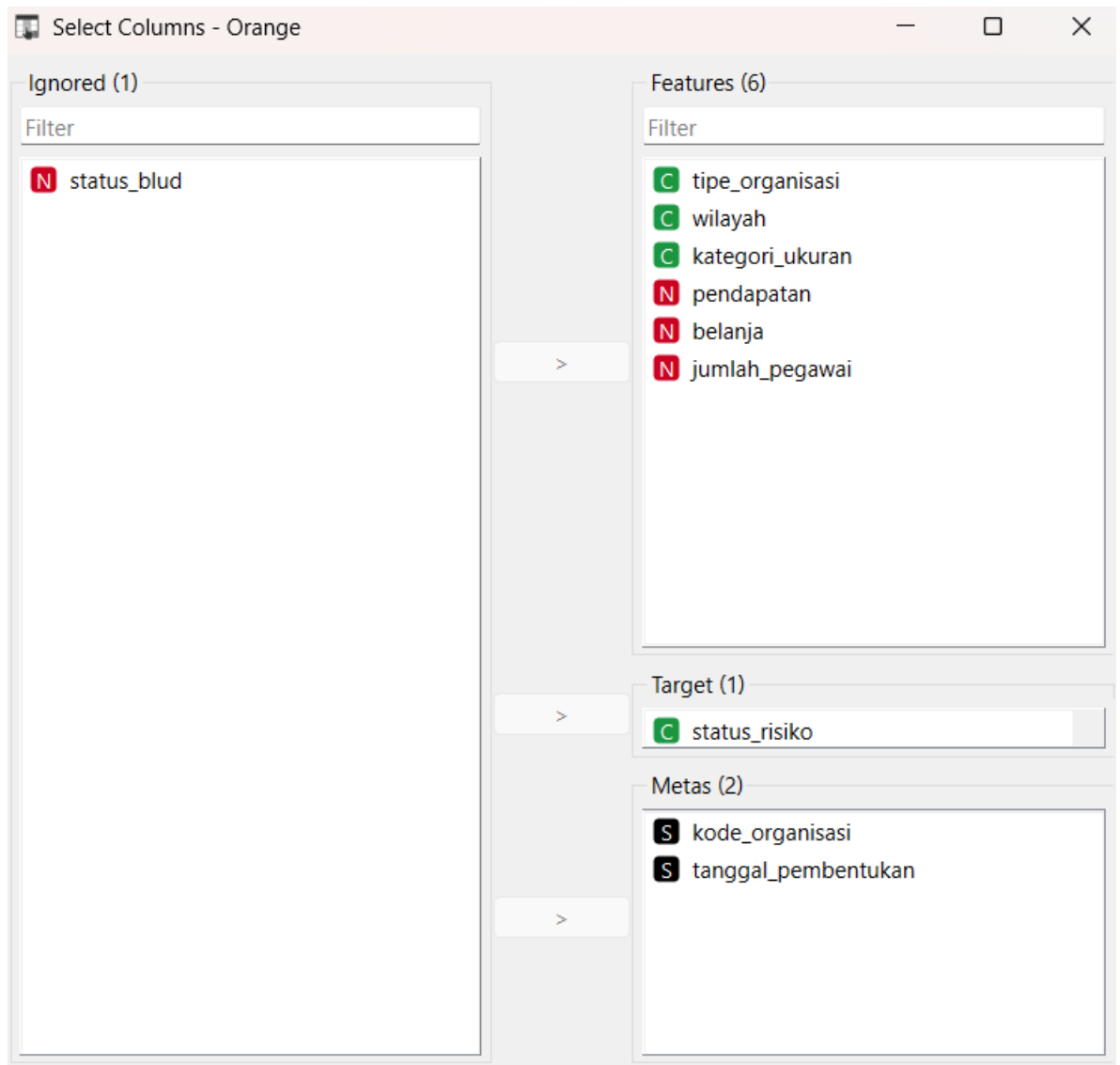


4. Encoding

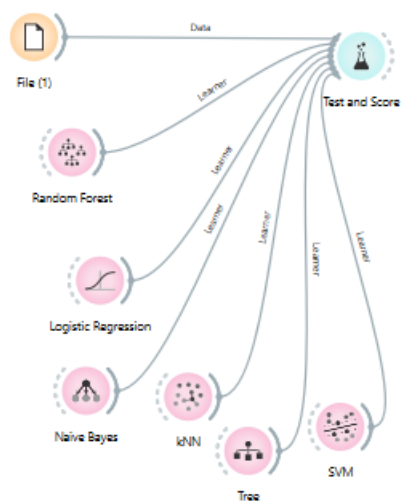
Atribut status_blud yang berisi True dan False diubah menjadi 1 dan 0



5. Memilih atribut yang digunakan sebagai feature dan class



C. Pembuatan Model



D. Evaluasi model

Test and Score - Orange

File Edit View Window Help

☒ Cross validation

Number of folds: 10

☒ Stratified

☐ Cross validation by feature

☐ Random sampling

Repeat train/test: 10

Training set size: 95 %

☒ Stratified

☐ Leave one out

☐ Test on train data

☐ Test on test data

Evaluation results for target (None, show average over classes)

Model	AUC	CA	F1	Prec	Recall	MCC
Random Forest	0.521	0.696	0.617	0.588	0.696	0.092
Tree	0.522	0.556	0.560	0.565	0.556	0.038
Logistic Regression	0.523	0.709	0.596	0.599	0.709	0.052
Naive Bayes	0.528	0.676	0.609	0.574	0.676	0.075
SVM	0.543	0.718	0.618	0.643	0.718	0.146
kNN	0.569	0.696	0.622	0.614	0.696	0.102

Berdasarkan evaluasi dengan test and score dan confusion matrix, ditentukan model terbaik adalah Tree karena nilai CA, F1 score, Precision, recall, dan MCC selalu paling atas (berarti Tree menunjukkan keakuratan, kemampuan diskriminasi antar target yang baik, dan keseimbangan antara precision dan recall), hanya AUC yang peringkat 2 di bawah random forest. Selain itu, dengan confusion matrix, Tree menunjukkan hasil yang paling mendekati actual.

E. Prediksi

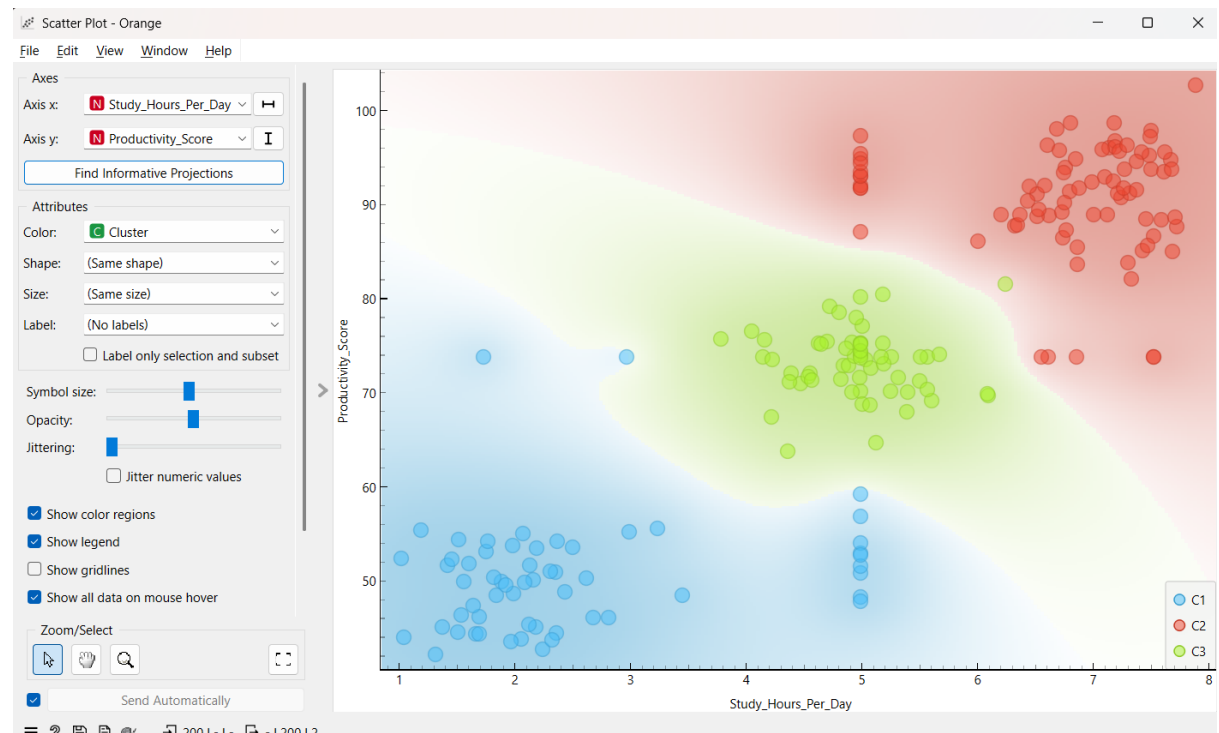
Jml data pada masing masing kelas hasil prediksi = 6

Pertanyaan

1. Sudah terjawab
2. Sudah terjawab
3. Preprocessing yang dilakukan adalah normalize, hal ini diperlukan untuk membuat konsistensi data dan mengurangi bias.
4. Sudah terjawab
5. Sudah terjawab
6. Interpretasi confusion matrix

Soal 2 klasterisasi

Algoritma yang memberikan hasil clustering terbaik adalah k means (dengan jumlah cluster ditentukan adalah 3) karena dapat memisahkan data dengan baik (3 kelompok data saling mengumpul dan terpisah), hanya ada sedikit data yang masuk ke cluster data lain. Hal ini berarti kemampuan k means untuk mendiskriminasi data cukup baik



Workflow

