
EvalML Documentation

Release 0.6.0

Feature Labs, Inc.

Dec 17, 2019

GETTING STARTED

1	Install	3
2	Quick Start	5
	Index	139



EvalML

EvalML is an AutoML library that builds, optimizes, and evaluates machine learning pipelines using domain-specific objective functions.

Combined with [Featuretools](#) and [Compose](#), EvalML can be used to create end-to-end machine learning solutions for classification and regression problems.

INSTALL

EvalML is available for Python 3.5+. It can be installed by running the following command:

```
pip install evaml --extra-index-url https://install.featurelabs.com/<license>/
```


QUICK START

```
[1]: import evalml
      from evalml import AutoClassificationSearch
```

2.1 Load Data

First, we load in the features and outcomes we want to use to train our model

```
[2]: X, y = evalml.demos.load_breast_cancer()
```

2.2 Configure search

EvalML has many options to configure the pipeline search. At the minimum, we need to define an objective function. For simplicity, we will use the F1 score in this example. However, the real power of EvalML is in using domain-specific *objective functions* or *building your own*.

Below EvalML utilizes Bayesian optimization (EvalML's default optimizer) to search and find the best pipeline defined by the given objective.

```
[3]: automl = AutoClassificationSearch(objective="f1",
                                     max_pipelines=5)
```

In order to validate the results of the pipeline creation and optimization process, we will save some of our data as a holdout set.

```
[4]: X_train, X_holdout, y_train, y_holdout = evalml.preprocessing.split_data(X, y, test_
      ↪size=.2)
```

When we call `.search()`, the search for the best pipeline will begin. There is no need to wrangle with missing data or categorical variables as EvalML includes various preprocessing steps (like imputation, one-hot encoding, feature selection) to ensure you're getting the best results. As long as your data is in a single table, EvalML can handle it. If not, you can reduce your data to a single table by utilizing [Featuretools](#) and its Entity Sets.

You can find more information on pipeline components and how to integrate your own custom pipelines into EvalML [here](#).

```
[5]: automl.search(X_train, y_train)
```

```
*****
* Beginning pipeline search *
*****
```

Optimizing for F1. Greater score is better.

Searching up to 5 pipelines.

Possible model types: xgboost, linear_model, random_forest

```
FigureWidget({
  'data': [{'mode': 'lines+markers',
            'name': 'Best Score',
            'type'...
```

XGBoost Classifier w/ One Hot Encod...	20%	Elapsed:00:02
XGBoost Classifier w/ One Hot Encod...	40%	Elapsed:00:04
Random Forest Classifier w/ One Hot...	60%	Elapsed:00:14
XGBoost Classifier w/ One Hot Encod...	80%	Elapsed:00:16
Logistic Regression Classifier w/ O...	100%	Elapsed:00:18
Optimization finished	100%	Elapsed:00:18

2.3 See Pipeline Rankings

After the search is finished we can view all of the pipelines searched, ranked by score. Internally, EvalML performs cross validation to score the pipelines. If it notices a high variance across cross validation folds, it will warn you. EvalML also provides additional *guardrails* to analyze your data to assist you in producing the best performing pipeline.

```
[6]: automl.rankings
```

```
[6]:   id      pipeline_class_name      score  high_variance_cv  \
0    4  LogisticRegressionPipeline  0.974067             False
1    2    RFClassificationPipeline  0.973981             False
2    1                XGBoostPipeline  0.961555             False
3    0                XGBoostPipeline  0.954991             False
4    3                XGBoostPipeline  0.951199             False
```

```
                                parameters
0  {'penalty': 'l2', 'C': 8.444214828324364, 'imp...
1  {'n_estimators': 569, 'max_depth': 22, 'impute...
2  {'eta': 0.38438170729269994, 'min_child_weight...
3  {'eta': 0.5928446182250184, 'min_child_weight'...
4  {'eta': 0.5288949197529046, 'min_child_weight'...
```

2.4 Describe pipeline

If we are interested in see more details about the pipeline, we can describe it using the `id` from the rankings table:

```
[7]: automl.describe_pipeline(3)
```

```
*****
* XGBoost Classifier w/ One Hot Encoder + Simple Imputer + RF Classifier Select From_
  ↳Model *
```

(continues on next page)

(continued from previous page)

```

*****
Problem Types: Binary Classification, Multiclass Classification
Model Type: XGBoost Classifier
Objective to Optimize: F1 (greater is better)
Number of features: 10

Pipeline Steps
=====
1. One Hot Encoder
2. Simple Imputer
   * impute_strategy : most_frequent
3. RF Classifier Select From Model
   * percent_features : 0.34402219881309576
   * threshold : -inf
4. XGBoost Classifier
   * eta : 0.5288949197529046
   * max_depth : 6
   * min_child_weight : 6.112401049845392

Training
=====
Training for Binary Classification problems.
Total training time (including CV): 2.4 seconds

Cross Validation
-----

```

	F1	Precision	Recall	AUC	Log Loss	MCC	# Training	# Testing
0	0.922	0.908	0.937	0.980	0.188	0.788	303.000	152.000
1	0.969	0.959	0.979	0.986	0.128	0.916	303.000	152.000
2	0.963	0.978	0.947	0.993	0.125	0.903	304.000	151.000
mean	0.951	0.948	0.954	0.986	0.147	0.869	-	-
std	0.025	0.036	0.022	0.007	0.036	0.070	-	-
coef of var	0.027	0.038	0.023	0.007	0.242	0.081	-	-

2.5 Select Best pipeline

We can now select best pipeline and score it on our holdout data:

```

[8]: pipeline = automl.best_pipeline
      pipeline.score(X_holdout, y_holdout)
[8]: (0.965034965034965, {})

```

We can also visualize the structure of our pipeline:

```

[9]: pipeline.plot()
[9]:

```

2.6 Whats next?

Head into the more in-depth automated walkthrough [here](#) or any advanced topics below.

2.6.1 Objective Functions

The **objective function** is what EvalML maximizes (or minimizes) as it completes the pipeline search. As it gets feedback from building pipelines, it tunes the hyperparameters to build optimized models. Therefore, it is critical to have an objective function that captures the how the model's predictions will be used in a business setting.

List of Available Objective Functions

Most AutoML libraries optimize for generic machine learning objective functions. Frequently, the scores produced by the generic machine learning objective diverge from how the model will be evaluated in the real world.

In EvalML, we can train and optimize the model for a specific problem by optimizing a domain-specific objectives functions or by defining our own custom objective function.

Currently, EvalML has two domain specific objective functions with more being developed. For more information on these objective functions click on the links below.

- [Fraud Detection](#)
- [Lead Scoring](#)

Build your own objective Functions

Often times, the objective function is very specific to the use-case or business problem. To get the right objective to optimize requires thinking through the decisions or actions that will be taken using the model and assigning the cost/benefit to doing that correctly or incorrectly based on known outcomes in the training data.

Once you have determined the objective for your business, you can provide that to EvalML to optimize by defining a custom objective function. Read more [here](#).

2.6.2 Building a Fraud Prediction Model with EvalML

In this demo, we will build an optimized fraud prediction model using EvalML. To optimize the pipeline, we will set up an objective function to minimize the percentage of total transaction value lost to fraud. At the end of this demo, we also show you how introducing the right objective during the training is over 4x better than using a generic machine learning metric like AUC.

```
[1]: import evalml
from evalml import AutoClassificationSearch
from evalml.objectives import FraudCost
```

Configure “Cost of Fraud”

To optimize the pipelines toward the specific business needs of this model, you can set your own assumptions for the cost of fraud. These parameters are

- `retry_percentage` - what percentage of customers will retry a transaction if it is declined?
- `interchange_fee` - how much of each successful transaction do you collect?
- `fraud_payout_percentage` - the percentage of fraud will you be unable to collect
- `amount_col` - the column in the data the represents the transaction amount

Using these parameters, EvalML determines attempt to build a pipeline that will minimize the financial loss due to fraud.

```
[2]: fraud_objective = FraudCost(retry_percentage=.5,
                                interchange_fee=.02,
                                fraud_payout_percentage=.75,
                                amount_col='amount')
```

Search for best pipeline

In order to validate the results of the pipeline creation and optimization process, we will save some of our data as a holdout set

```
[3]: X, y = evalml.demos.load_fraud(n_rows=2500)
```

```

      Number of Features
Boolean                1
Categorical            6
Numeric               5

Number of training examples: 2500
Labels
False    85.92%
True     14.08%
Name: fraud, dtype: object
```

EvalML natively supports one-hot encoding. Here we keep 1 out of the 6 categorical columns to decrease computation time.

```
[4]: X = X.drop(['datetime', 'expiration_date', 'country', 'region', 'provider'], axis=1)

X_train, X_holdout, y_train, y_holdout = evalml.preprocessing.split_data(X, y, test_
↳size=0.2, random_state=0)

print(X.dtypes)

card_id          int64
store_id         int64
amount           int64
currency         object
customer_present bool
lat              float64
lng              float64
dtype: object
```

Because the fraud labels are binary, we will use `AutoClassificationSearch`. When we call `.search()`, the search for the best pipeline will begin.

```
[5]: automl = AutoClassificationSearch(objective=fraud_objective,
                                     additional_objectives=['auc', 'recall', 'precision
↳'],
                                     max_pipelines=5)

automl.search(X_train, y_train)

*****
* Beginning pipeline search *
*****

Optimizing for Fraud Cost. Lower score is better.
```

(continues on next page)

(continued from previous page)

```
Searching up to 5 pipelines.
Possible model types: xgboost, linear_model, random_forest

FigureWidget({
  'data': [{'mode': 'lines+markers',
            'name': 'Best Score',
            'type'...

XGBoost Classifier w/ One Hot Encod... 20%|          | Elapsed:00:06
XGBoost Classifier w/ One Hot Encod... 40%|          | Elapsed:00:11
Random Forest Classifier w/ One Hot... 60%|          | Elapsed:00:28
XGBoost Classifier w/ One Hot Encod... 80%|          | Elapsed:00:34
Logistic Regression Classifier w/ O... 100%|         | Elapsed:00:38
Optimization finished                  100%|         | Elapsed:00:38
```

View rankings and select pipeline

Once the fitting process is done, we can see all of the pipelines that were searched, ranked by their score on the fraud detection objective we defined

```
[6]: automl.rankings

[6]:   id  pipeline_class_name  score  high_variance_cv  \
0    2  RFClassificationPipeline  0.007770          False
1    0      XGBoostPipeline  0.007838          False
2    1      XGBoostPipeline  0.007838          False
3    3      XGBoostPipeline  0.007838          False
4    4  LogisticRegressionPipeline  0.008914          False

                                parameters
0  {'n_estimators': 569, 'max_depth': 22, 'impute...
1  {'eta': 0.5928446182250184, 'min_child_weight'...
2  {'eta': 0.38438170729269994, 'min_child_weight'...
3  {'eta': 0.5288949197529046, 'min_child_weight'...
4  {'penalty': 'l2', 'C': 8.444214828324364, 'imp...
```

to select the best pipeline we can run

```
[7]: best_pipeline = automl.best_pipeline
```

Describe pipeline

You can get more details about any pipeline. Including how it performed on other objective functions.

```
[8]: automl.describe_pipeline(automl.rankings.iloc[0]["id"])

*****
* Random Forest Classifier w/ One Hot Encoder + Simple Imputer + RF Classifier Select_
↳ From Model *
*****

Problem Types: Binary Classification, Multiclass Classification
Model Type: Random Forest
```

(continues on next page)

(continued from previous page)

Objective to Optimize: Fraud Cost (lower is better)
 Number of features: 6

Pipeline Steps

=====

1. One Hot Encoder
2. Simple Imputer
 - * impute_strategy : most_frequent
3. RF Classifier Select From Model
 - * percent_features : 0.8593661614465293
 - * threshold : -inf
4. Random Forest Classifier
 - * n_estimators : 569
 - * max_depth : 22

Training

=====

Training for Binary Classification problems.
 Total training time (including CV): 16.1 seconds

Cross Validation

	Fraud Cost	AUC	Recall	Precision	# Training	# Testing
0	0.008	0.856	1.000	0.143	1333.000	667.000
1	0.008	0.880	1.000	0.142	1333.000	667.000
2	0.008	0.846	1.000	0.142	1334.000	666.000
mean	0.008	0.861	1.000	0.142	-	-
std	0.000	0.018	0.000	0.001	-	-
coef of var	0.003	0.020	0.000	0.005	-	-

Evaluate on hold out

Finally, we retrain the best pipeline on all of the training data and evaluate on the holdout

```
[9]: best_pipeline.fit(X_train, y_train)
[9]: <evalml.pipelines.classification.random_forest.RFClassificationPipeline at_
     ↪ 0x7fc5b6ce44e0>
```

Now, we can score the pipeline on the hold out data using both the fraud cost score and the AUC.

```
[10]: best_pipeline.score(X_holdout, y_holdout, other_objectives=["auc", fraud_objective])
[10]: (0.007738666614701976,
      OrderedDict([('AUC', 0.8219767441860466),
                    ('Fraud Cost', 0.007738666614701976)]))
```

Why optimize for a problem-specific objective?

To demonstrate the importance of optimizing for the right objective, let's search for another pipeline using AUC, a common machine learning metric. After that, we will score the holdout data using the fraud cost objective to see how the best pipelines compare.

```
[11]: automl_auc = AutoClassificationSearch(objective='auc',
                                          additional_objectives=['recall', 'precision'],
                                          max_pipelines=5)

automl_auc.search(X_train, y_train)

*****
* Beginning pipeline search *
*****

Optimizing for AUC. Greater score is better.

Searching up to 5 pipelines.
Possible model types: xgboost, linear_model, random_forest

FigureWidget({
  'data': [{'mode': 'lines+markers',
            'name': 'Best Score',
            'type'...

XGBoost Classifier w/ One Hot Encod... 20%|          | Elapsed:00:03
XGBoost Classifier w/ One Hot Encod... 40%|          | Elapsed:00:06
Random Forest Classifier w/ One Hot... 60%|          | Elapsed:00:21
XGBoost Classifier w/ One Hot Encod... 80%|          | Elapsed:00:25
Logistic Regression Classifier w/ O... 100%|         | Elapsed:00:26
Optimization finished                  100%|         | Elapsed:00:26
```

like before, we can look at the rankings and pick the best pipeline

```
[12]: automl_auc.rankings

[12]:   id      pipeline_class_name      score  high_variance_cv  \
0    2      RFClassificationPipeline  0.868852             False
1    1              XGBoostPipeline  0.852802             False
2    3              XGBoostPipeline  0.846817             False
3    0              XGBoostPipeline  0.843807             False
4    4  LogisticRegressionPipeline  0.677654             False

                                parameters
0  {'n_estimators': 569, 'max_depth': 22, 'impute...
1  {'eta': 0.38438170729269994, 'min_child_weight...
2  {'eta': 0.5288949197529046, 'min_child_weight'...
3  {'eta': 0.5928446182250184, 'min_child_weight'...
4  {'penalty': 'l2', 'C': 8.444214828324364, 'imp...
```

```
[13]: best_pipeline_auc = automl_auc.best_pipeline

# train on the full training data
best_pipeline_auc.fit(X_train, y_train)

[13]: <evalml.pipelines.classification.random_forest.RFClassificationPipeline at_
      ↪0x7fc5b61064a8>
```

```
[14]: # get the fraud score on holdout data
best_pipeline_auc.score(X_holdout, y_holdout, other_objectives=["auc", fraud_
      ↪objective])
```



```
[14]: (0.8354983388704318,
      OrderedDict([('AUC', 0.8354983388704318),
                    ('Fraud Cost', 0.03655681280302016)]))

[15]: # fraud score on fraud optimized again
best_pipeline.score(X_holdout, y_holdout, other_objectives=["auc", fraud_objective])

[15]: (0.007738666614701976,
      OrderedDict([('AUC', 0.8219767441860466),
                    ('Fraud Cost', 0.007738666614701976)]))
```

When we optimize for AUC, we can see that the AUC score from this pipeline is better than the AUC score from the pipeline optimized for fraud cost. However, the losses due to fraud are over 3% of the total transaction amount when optimized for AUC and under 1% when optimized for fraud cost. As a result, we lose more than 2% of the total transaction amount by not optimizing for fraud cost specifically.

This happens because optimizing for AUC does not take into account the user-specified `retry_percentage`, `interchange_fee`, `fraud_payout_percentage` values. Thus, the best pipelines may produce the highest AUC but may not actually reduce the amount loss due to your specific type fraud.

This example highlights how performance in the real world can diverge greatly from machine learning metrics.

2.6.3 Building a Lead Scoring Model with EvalML

In this demo, we will build an optimized lead scoring model using EvalML. To optimize the pipeline, we will set up an objective function to maximize the revenue generated with true positives while taking into account the cost of false positives. At the end of this demo, we also show you how introducing the right objective during the training is over 6x better than using a generic machine learning metric like AUC.

```
[1]: import evalml
      from evalml import AutoClassificationSearch
      from evalml.objectives import LeadScoring
```

Configure LeadScoring

To optimize the pipelines toward the specific business needs of this model, you can set your own assumptions for how much value is gained through true positives and the cost associated with false positives. These parameters are

- `true_positive` - dollar amount to be gained with a successful lead
- `false_positive` - dollar amount to be lost with an unsuccessful lead

Using these parameters, EvalML builds a pipeline that will maximize the amount of revenue per lead generated.

```
[2]: lead_scoring_objective = LeadScoring(
      true_positives=1000,
      false_positives=-10
    )
```

Dataset

We will be utilizing a dataset detailing a customer's job, country, state, zip, online action, the dollar amount of that action and whether they were a successful lead.

```
[3]: import pandas as pd

customers = pd.read_csv('s3://featurelabs-static/lead_scoring_ml_apps/customers.csv')
interactions = pd.read_csv('s3://featurelabs-static/lead_scoring_ml_apps/interactions.
↪csv')
leads = pd.read_csv('s3://featurelabs-static/lead_scoring_ml_apps/previous_leads.csv')

X = customers.merge(interactions, on='customer_id').merge(leads, on='customer_id')
y = X['label']

X = X.drop(['customer_id', 'date_registered', 'birthday', 'phone', 'email',
           'owner', 'company', 'id', 'time_x',
           'session', 'referrer', 'time_y', 'label'], axis=1)

display(X.head())
```

	job	country	state	zip	action	amount
0	Engineer, mining	NaN	NY	60091.0	page_view	NaN
1	Psychologist, forensic	US	CA	NaN	purchase	135.23
2	Psychologist, forensic	US	CA	NaN	page_view	NaN
3	Air cabin crew	US	NaN	60091.0	download	NaN
4	Air cabin crew	US	NaN	60091.0	page_view	NaN

Search for best pipeline

In order to validate the results of the pipeline creation and optimization process, we will save some of our data as a holdout set

EvalML natively supports one-hot encoding and imputation so the above NaN and categorical values will be taken care of.

```
[4]: X_train, X_holdout, y_train, y_holdout = evalml.preprocessing.split_data(X, y, test_
↪size=0.2, random_state=0)

print(X.dtypes)

job          object
country      object
state        object
zip          float64
action       object
amount       float64
dtype: object
```

Because the lead scoring labels are binary, we will use `AutoClassificationSearch`. When we call `.search()`, the search for the best pipeline will begin.

```
[5]: automl = AutoClassificationSearch(objective=lead_scoring_objective,
                                     additional_objectives=['auc'],
                                     max_pipelines=5)

automl.search(X_train, y_train)

*****
* Beginning pipeline search *
*****
```

(continues on next page)

(continued from previous page)

```
Optimizing for Lead Scoring. Greater score is better.
```

```
Searching up to 5 pipelines.
```

```
Possible model types: random_forest, linear_model, xgboost
```

```
FigureWidget({
  'data': [{'mode': 'lines+markers',
            'name': 'Best Score',
            'type'...
```

```
XGBoost Classifier w/ One Hot Encod... 20%|          | Elapsed:00:06
XGBoost Classifier w/ One Hot Encod... 40%|          | Elapsed:00:13
Random Forest Classifier w/ One Hot... 60%|          | Elapsed:00:31
XGBoost Classifier w/ One Hot Encod... 80%|          | Elapsed:00:38
Logistic Regression Classifier w/ O... 100%|         | Elapsed:00:44
Optimization finished                  100%|         | Elapsed:00:44
```

View rankings and select pipeline

Once the fitting process is done, we can see all of the pipelines that were searched, ranked by their score on the lead scoring objective we defined

```
[6]: automl.rankings
```

```
[6]:   id      pipeline_class_name      score  high_variance_cv  \
0    0          XGBoostPipeline  15.997332             False
1    2  RFClassificationPipeline  15.095675             False
2    1          XGBoostPipeline  14.474133             False
3    3          XGBoostPipeline  13.893124             False
4    4  LogisticRegressionPipeline  13.436065              True

                                parameters
0  {'eta': 0.59284446182250184, 'min_child_weight'...
1  {'n_estimators': 569, 'max_depth': 22, 'impute...
2  {'eta': 0.38438170729269994, 'min_child_weight'...
3  {'eta': 0.5288949197529046, 'min_child_weight'...
4  {'penalty': 'l2', 'C': 8.444214828324364, 'imp...
```

to select the best pipeline we can run

```
[7]: best_pipeline = automl.best_pipeline
```

Describe pipeline

You can get more details about any pipeline. Including how it performed on other objective functions.

```
[8]: automl.describe_pipeline(automl.rankings.iloc[0]["id"])
```

```
*****
* XGBoost Classifier w/ One Hot Encoder + Simple Imputer + RF Classifier Select From_
  ↳Model *
*****

Problem Types: Binary Classification, Multiclass Classification
```

(continues on next page)

(continued from previous page)

```

Model Type: XGBoost Classifier
Objective to Optimize: Lead Scoring (greater is better)
Number of features: 3

Pipeline Steps
=====
1. One Hot Encoder
2. Simple Imputer
   * impute_strategy : most_frequent
3. RF Classifier Select From Model
   * percent_features : 0.6273280598181127
   * threshold : -inf
4. XGBoost Classifier
   * eta : 0.5928446182250184
   * max_depth : 4
   * min_child_weight : 8.598391737229157

Training
=====
Training for Binary Classification problems.
Total training time (including CV): 6.4 seconds

Cross Validation
-----

```

	Lead Scoring	AUC	# Training	# Testing
0	13.774	0.500	3099.000	1550.000
1	16.626	0.500	3099.000	1550.000
2	17.592	0.500	3100.000	1549.000
mean	15.997	0.500	-	-
std	1.985	0.000	-	-
coef of var	0.124	0.000	-	-

Evaluate on hold out

Finally, we retrain the best pipeline on all of the training data and evaluate on the holdout

```

[9]: best_pipeline.fit(X_train, y_train)
[9]: <evalml.pipelines.classification.xgboost.XGBoostPipeline at 0x7efefb2d2fd0>

```

Now, we can score the pipeline on the hold out data using both the lead scoring score and the AUC.

```

[10]: best_pipeline.score(X_holdout, y_holdout, other_objectives=["auc", lead_scoring_
    ↪objective])
[10]: (15.382631126397248,
    OrderedDict([('AUC', 0.5), ('Lead Scoring', 15.382631126397248)]))

```

Why optimize for a problem-specific objective?

To demonstrate the importance of optimizing for the right objective, let's search for another pipeline using AUC, a common machine learning metric. After that, we will score the holdout data using the lead scoring objective to see how the best pipelines compare.

```
[11]: automl_auc = evalml.AutoClassificationSearch(objective='auc',
                                                additional_objectives=[],
                                                max_pipelines=5)

automl_auc.search(X_train, y_train)

*****
* Beginning pipeline search *
*****

Optimizing for AUC. Greater score is better.

Searching up to 5 pipelines.
Possible model types: random_forest, linear_model, xgboost

FigureWidget({
  'data': [{'mode': 'lines+markers',
            'name': 'Best Score',
            'type'...

XGBoost Classifier w/ One Hot Encod... 20%|          | Elapsed:00:02
XGBoost Classifier w/ One Hot Encod... 40%|          | Elapsed:00:05
Random Forest Classifier w/ One Hot... 60%|          | Elapsed:00:20
XGBoost Classifier w/ One Hot Encod... 80%|          | Elapsed:00:22
Logistic Regression Classifier w/ O... 100%|         | Elapsed:00:24
Optimization finished                  100%|         | Elapsed:00:24
```

like before, we can look at the rankings and pick the best pipeline

```
[12]: automl_auc.rankings

[12]:   id      pipeline_class_name      score  high_variance_cv  \
0    4  LogisticRegressionPipeline  0.921954                False
1    1                XGBoostPipeline  0.556359                False
2    2      RFClassificationPipeline  0.554214                False
3    0                XGBoostPipeline  0.537619                False
4    3                XGBoostPipeline  0.500000                False

                                parameters
0  {'penalty': 'l2', 'C': 8.444214828324364, 'imp...
1  {'eta': 0.38438170729269994, 'min_child_weight...
2  {'n_estimators': 569, 'max_depth': 22, 'impute...
3  {'eta': 0.5928446182250184, 'min_child_weight'...
4  {'eta': 0.5288949197529046, 'min_child_weight'...
```

```
[13]: best_pipeline_auc = automl_auc.best_pipeline

# train on the full training data
best_pipeline_auc.fit(X_train, y_train)

[13]: <evalml.pipelines.classification.logistic_regression.LogisticRegressionPipeline at_
      ↪0x7efef99a7128>
```

```
[14]: # get the auc and lead scoring score on holdout data
best_pipeline_auc.score(X_holdout, y_holdout, other_objectives=["auc", lead_scoring_
      ↪objective])
```

```
[14]: (0.9272061045633122,
OrderedDict([('AUC', 0.9272061045633122),
              ('Lead Scoring', -0.017196904557179708)]))
```

When we optimize for AUC, we can see that the AUC score from this pipeline is better than the AUC score from the pipeline optimized for lead scoring. However, the revenue per lead gained was only \$7 per lead when optimized for AUC and was \$45 when optimized for lead scoring. As a result, we would gain up to 6x the amount of revenue if we optimized for lead scoring.

This happens because optimizing for AUC does not take into account the user-specified `true_positive` (dollar amount to be gained with a successful lead) and `false_positive` (dollar amount to be lost with an unsuccessful lead) values. Thus, the best pipelines may produce the highest AUC but may not actually generate the most revenue through lead scoring.

This example highlights how performance in the real world can diverge greatly from machine learning metrics.

2.6.4 Custom Objective Functions

Often times, the objective function is very specific to the use-case or business problem. To get the right objective to optimize requires thinking through the decisions or actions that will be taken using the model and assigning a cost/benefit to doing that correctly or incorrectly based on known outcomes in the training data.

Once you have determined the objective for your business, you can provide that to EvalML to optimize by defining a custom objective function.

How to Create a Objective Function

To create a custom objective function, we must define 2 functions

- The **“objective function”**: this function takes the predictions, true labels, and any other information about the future and returns a score of how well the model performed.
- The **“decision function”**: this function takes prediction probabilities that were output from the model and a threshold and returns a prediction.

To evaluate a particular model, EvalML automatically finds the best threshold to pass to the decision function to generate predictions and then scores the resulting predictions using the objective function. The score from the objective function determines which set of pipeline hyperparameters EvalML will try next.

To give a concrete example, let’s look at how the fraud detection objective function is built.

```
[1]: from evalml.objectives.objective_base import ObjectiveBase

class FraudCost(ObjectiveBase):
    """Score the percentage of money lost of the total transaction amount process due_
    to fraud"""
    name = "Fraud Cost"
    needs_fitting = True
    greater_is_better = False
    uses_extra_columns = True
    fit_needs_proba = True
    score_needs_proba = False

    def __init__(self, retry_percentage=.5, interchange_fee=.02,
                 fraud_payout_percentage=1.0, amount_col='amount', verbose=False):
        """Create instance of FraudCost
```

(continues on next page)

(continued from previous page)

```

    Args:
        retry_percentage (float): what percentage of customers will retry a
        ↪ transaction if it
            is declined? Between 0 and 1. Defaults to .5

        interchange_fee (float): how much of each successful transaction do you
        ↪ collect?
            Between 0 and 1. Defaults to .02

        fraud_payout_percentage (float): how percentage of fraud will you be
        ↪ unable to collect.
            Between 0 and 1. Defaults to 1.0

        amount_col (str): name of column in data that contains the amount.
        ↪ defaults to "amount"
        """
        self.retry_percentage = retry_percentage
        self.interchange_fee = interchange_fee
        self.fraud_payout_percentage = fraud_payout_percentage
        self.amount_col = amount_col
        super().__init__(verbose=verbose)

    def decision_function(self, y_predicted, extra_cols, threshold):
        """Determine if transaction is fraud given predicted probabilities,
        dataframe with transaction amount, and threshold"""

        transformed_probs = (y_predicted * extra_cols[self.amount_col])
        return transformed_probs > threshold

    def objective_function(self, y_predicted, y_true, extra_cols):
        ↪ dataframe
        """Calculate amount lost to fraud given predictions, true values, and
            with transaction amount"""

        # extract transaction using the amount columns in users data
        transaction_amount = extra_cols[self.amount_col]

        # amount paid if transaction is fraud
        fraud_cost = transaction_amount * self.fraud_payout_percentage

        # money made from interchange fees on transaction
        interchange_cost = transaction_amount * (1 - self.retry_percentage) * self.
        ↪ interchange_fee

        # calculate cost of missing fraudulent transactions
        false_negatives = (y_true & ~y_predicted) * fraud_cost

        # calculate money lost from fees
        false_positives = (~y_true & y_predicted) * interchange_cost

        loss = false_negatives.sum() + false_positives.sum()

        loss_per_total_processed = loss / transaction_amount.sum()

        return loss_per_total_processed

```

2.6.5 Setting up pipeline search

Designing the right machine learning pipeline and picking the best parameters is a time-consuming process that relies on a mix of data science intuition as well as trial and error. EvalML streamlines the process of selecting the best modeling algorithms and parameters, so data scientists can focus their energy where it is most needed.

How it works

EvalML selects and tunes machine learning pipelines built of numerous steps. This includes encoding categorical data, missing value imputation, feature selection, feature scaling, and finally machine learning. As EvalML tunes pipelines, it uses the objective function selected and configured by the user to guide its search.

At each iteration, EvalML uses cross-validation to generate an estimate of the pipeline's performances. If a pipeline has high variance across cross-validation folds, it will provide a warning. In this case, the pipeline may not perform reliably in the future.

EvalML is designed to work well out of the box. However, it provides numerous methods for you to control the search described below.

Selecting problem type

EvalML supports both classification and regression problems. You select your problem type by importing the appropriate class.

```
[1]: import evalml
      from evalml import AutoClassificationSearch, AutoRegressionSearch

[2]: AutoClassificationSearch()
[2]: <evalml.automl.auto_classification_search.AutoClassificationSearch at 0x7f98c94fd3c8>

[3]: AutoRegressionSearch()
[3]: <evalml.automl.auto_regression_search.AutoRegressionSearch at 0x7f98c7cbcc50>
```

Setting the Objective Function

The only required parameter to start searching for pipelines is the objective function. Most domain-specific objective functions require you to specify parameters based on your business assumptions. You can do this before you initialize your pipeline search. For example

```
[4]: from evalml.objectives import FraudCost

      fraud_objective = FraudCost(
          retry_percentage=.5,
          interchange_fee=.02,
          fraud_payout_percentage=.75,
          amount_col='amount'
      )

      AutoClassificationSearch(objective=fraud_objective)

[4]: <evalml.automl.auto_classification_search.AutoClassificationSearch at 0x7f98c7cd87b8>
```


Evaluate on Additional Objectives

Additional objectives can be scored on during the evaluation process. To add another objective, use the `additional_objectives` parameter in `AutoClassificationSearch` or `AutoRegressionSearch`. The results of these additional objectives will then appear in the results of `describe_pipeline`.

```
[5]: from evalml.objectives import FraudCost

fraud_objective = FraudCost(
    retry_percentage=.5,
    interchange_fee=.02,
    fraud_payout_percentage=.75,
    amount_col='amount'
)

AutoClassificationSearch(objective='AUC', additional_objectives=[fraud_objective])

[5]: <evalml.automl.auto_classification_search.AutoClassificationSearch at 0x7f98c7c644e0>
```

Selecting Model Types

By default, all model types are considered. You can control which model types to search with the `model_types` parameters

```
[6]: automl = AutoClassificationSearch(objective="f1",
                                     model_types=["random_forest"])
```

you can see the possible pipelines that will be searched after initialization

```
[7]: automl.possible_pipelines

[7]: [evalml.pipelines.classification.random_forest.RFClassificationPipeline]
```

you can see a list of all supported models like this

```
[8]: evalml.list_model_types("binary") # `binary` for binary classification and
    ↪ `multiclass` for multiclass classification

[8]: [<ModelTypes.LINEAR_MODEL: 'linear_model'>,
      <ModelTypes.XGBOOST: 'xgboost'>,
      <ModelTypes.RANDOM_FOREST: 'random_forest'>]

[9]: evalml.list_model_types("regression")

[9]: [<ModelTypes.LINEAR_MODEL: 'linear_model'>,
      <ModelTypes.RANDOM_FOREST: 'random_forest'>]
```

Limiting Search Time

You can limit the search time by specifying a maximum number of pipelines and/or a maximum amount of time. EvalML won't build new pipelines after the maximum time has passed or the maximum number of pipelines have been built. If a limit is not set, then a maximum of 5 pipelines will be built.

The maximum search time can be specified as an integer in seconds or as a string in seconds, minutes, or hours.

```
[10]: AutoClassificationSearch(objective="f1",
                             max_pipelines=5,
                             max_time=60)

AutoClassificationSearch(objective="f1",
                          max_time="1 minute")

[10]: <evalml automl auto_classification_search.AutoClassificationSearch at 0x7f98c7cdc7f0>
```

To start, EvalML samples 10 sets of hyperparameters chosen randomly for each possible pipeline. Therefore, we recommend setting `max_pipelines` at least 10 times the number of possible pipelines.

```
[11]: n_possible_pipelines = len(AutoClassificationSearch(objective="f1").possible_
    ↪pipelines)

[12]: AutoClassificationSearch(objective="f1",
                             max_time=60)

[12]: <evalml automl auto_classification_search.AutoClassificationSearch at 0x7f98c7c7df98>
```

Early Stopping

You can also limit search time by providing a patience value for early stopping. With a patience value, EvalML will stop searching when the best objective score has not been improved upon for `n` iterations. The patience value must be a positive integer. You can also provide a tolerance value where EvalML will only consider a score as an improvement over the best score if the difference was greater than the tolerance percentage.

```
[13]: from evalml.demos import load_diabetes

X, y = load_diabetes()
automl = AutoRegressionSearch(objective="MSE", patience=2, tolerance=0.01, max_
    ↪pipelines=10)
automl.search(X, y)

*****
* Beginning pipeline search *
*****

Optimizing for MSE. Lower score is better.

Searching up to 10 pipelines.
Possible model types: linear_model, random_forest

FigureWidget({
  'data': [{'mode': 'lines+markers',
            'name': 'Best Score',
            'type': ...

Random Forest Regressor w/ One Hot ... 10%|          | Elapsed:00:10
Random Forest Regressor w/ One Hot ... 20%|          | Elapsed:00:16
Linear Regressor w/ One Hot Encoder... 30%|          | Elapsed:00:16
Random Forest Regressor w/ One Hot ... 40%|          | Elapsed:00:25
Random Forest Regressor w/ One Hot ... 50%|          | Elapsed:00:36

2 iterations without improvement. Stopping search early...
Optimization finished                    50%|          | Elapsed:00:36
```

```
[14]: automl.rankings
```

	id	pipeline_class_name	score	high_variance_cv	\
0	2	LinearRegressionPipeline	3027.144520	False	
1	0	RFRegressionPipeline	3413.567159	False	
2	4	RFRegressionPipeline	3602.485675	False	
3	3	RFRegressionPipeline	3648.302977	False	
4	1	RFRegressionPipeline	3656.887580	False	


```

                                parameters
0  {'impute_strategy': 'mean', 'normalize': True,...
1  {'n_estimators': 569, 'max_depth': 22, 'impute...
2  {'n_estimators': 715, 'max_depth': 7, 'impute...
3  {'n_estimators': 609, 'max_depth': 7, 'impute...
4  {'n_estimators': 369, 'max_depth': 10, 'impute...

```

Control Cross Validation

EvalML cross-validates each model it tests during its search. By default it uses 3-fold cross-validation. You can optionally provide your own cross-validation method.

```
[15]: from sklearn.model_selection import StratifiedKFold

automl = AutoClassificationSearch(objective="f1",
                                cv=StratifiedKFold(5))
```

2.6.6 Exploring search results

After finishing a pipeline search, we can inspect the results. First, let's build a search of 10 different pipelines to explore.

```
[1]: import evalml
from evalml import AutoClassificationSearch

X, y = evalml.demos.load_breast_cancer()

automl = AutoClassificationSearch(objective="f1",
                                max_pipelines=5)

automl.search(X, y)
```

```

*****
* Beginning pipeline search *
*****

Optimizing for F1. Greater score is better.

Searching up to 5 pipelines.
Possible model types: random_forest, xgboost, linear_model

FigureWidget({
  'data': [{'mode': 'lines+markers',
            'name': 'Best Score',
            'type': ...

```

```

XGBoost Classifier w/ One Hot Encod... 20%|          | Elapsed:00:02
XGBoost Classifier w/ One Hot Encod... 40%|          | Elapsed:00:04
Random Forest Classifier w/ One Hot... 60%|          | Elapsed:00:14
XGBoost Classifier w/ One Hot Encod... 80%|          | Elapsed:00:17
Logistic Regression Classifier w/ O... 100%| Elapsed:00:18
Optimization finished                  100%| Elapsed:00:18

```

View Rankings

A summary of all the pipelines built can be returned as a pandas DataFrame. It is sorted by score. EvalML knows based on our objective function whether higher or lower is better.

```

[2]: automl.rankings

[2]:   id  pipeline_class_name  score  high_variance_cv  \
0    2  RFClassificationPipeline  0.973477          False
1    4  LogisticRegressionPipeline  0.971949          False
2    1  XGBoostPipeline  0.964896          False
3    0  XGBoostPipeline  0.959401          False
4    3  XGBoostPipeline  0.948566          False

                                parameters
0  {'n_estimators': 569, 'max_depth': 22, 'impute...
1  {'penalty': 'l2', 'C': 8.444214828324364, 'imp...
2  {'eta': 0.38438170729269994, 'min_child_weight...
3  {'eta': 0.5928446182250184, 'min_child_weight'...
4  {'eta': 0.5288949197529046, 'min_child_weight'...

```

Describe Pipeline

Each pipeline is given an `id`. We can get more information about any particular pipeline using that `id`. Here, we will get more information about the pipeline with `id = 0`.

```

[3]: automl.describe_pipeline(0)

*****
* XGBoost Classifier w/ One Hot Encoder + Simple Imputer + RF Classifier Select From_
↪Model *
*****

Problem Types: Binary Classification, Multiclass Classification
Model Type: XGBoost Classifier
Objective to Optimize: F1 (greater is better)
Number of features: 18

Pipeline Steps
=====
1. One Hot Encoder
2. Simple Imputer
   * impute_strategy : most_frequent
3. RF Classifier Select From Model
   * percent_features : 0.6273280598181127
   * threshold : -inf
4. XGBoost Classifier
   * eta : 0.5928446182250184

```

(continues on next page)

(continued from previous page)

```

* max_depth : 4
* min_child_weight : 8.598391737229157

Training
=====
Training for Binary Classification problems.
Total training time (including CV): 2.4 seconds

Cross Validation
-----

```

	F1	Precision	Recall	AUC	Log Loss	MCC	# Training	# Testing
0	0.979	0.975	0.983	0.998	0.081	0.944	379.000	190.000
1	0.954	0.950	0.958	0.984	0.142	0.876	379.000	190.000
2	0.945	0.949	0.941	0.981	0.159	0.853	380.000	189.000
mean	0.959	0.958	0.961	0.988	0.127	0.891	-	-
std	0.018	0.015	0.021	0.009	0.041	0.047	-	-
coef of var	0.018	0.015	0.022	0.009	0.324	0.053	-	-

Get Pipeline

We can get the object of any pipeline via their id as well:

```

[4]: automl.get_pipeline(0)
[4]: <evalml.pipelines.classification.xgboost.XGBoostPipeline at 0x7f1ab6b0be48>

```

Get best pipeline

If we specifically want to get the best pipeline, there is a convenient access

```

[5]: automl.best_pipeline
[5]: <evalml.pipelines.classification.random_forest.RFClassificationPipeline at 0x7f1ab62cfb00>

```

Feature Importances

We can get the feature importances of the resulting pipeline

```

[6]: pipeline = automl.get_pipeline(0)
      pipeline.feature_importances
[6]:

```

	feature	importance
0	worst perimeter	0.341811
1	worst radius	0.184930
2	mean concave points	0.163518
3	worst concave points	0.115095
4	mean concavity	0.047942
5	worst area	0.038873
6	worst concavity	0.032179
7	area error	0.028544
8	worst texture	0.022472
9	worst symmetry	0.015158

(continues on next page)

(continued from previous page)

```

10         mean smoothness      0.005401
11         radius error         0.004078
12         mean radius          0.000000
13         mean perimeter       0.000000
14         mean area            0.000000
15         mean compactness     0.000000
16         worst compactness    0.000000
17         worst fractal dimension 0.000000

```

We can also create a bar plot of the feature importances

```
[7]: pipeline.plot.feature_importances()
```

Data type cannot be displayed: application/vnd.plotly.v1+json, text/html

Plot ROC

For binary classification tasks, we can also plot the ROC plot of a specific pipeline:

```
[8]: automl.plot.generate_roc_plot(0)
```

Data type cannot be displayed: application/vnd.plotly.v1+json, text/html

Access raw results

You can also get access to all the underlying data like this

```
[9]: automl.results
```

```

[9]: {'pipeline_results': {0: {'id': 0,
    'pipeline_class_name': 'XGBoostPipeline',
    'pipeline_name': 'XGBoost Classifier w/ One Hot Encoder + Simple Imputer + RF_
    ↳Classifier Select From Model',
    'parameters': {'eta': 0.5928446182250184,
    'min_child_weight': 8.598391737229157,
    'max_depth': 4,
    'impute_strategy': 'most_frequent',
    'percent_features': 0.6273280598181127},
    'score': 0.9594006908767779,
    'high_variance_cv': False,
    'training_time': 2.430683135986328,
    'cv_data': [{'all_objective_scores': OrderedDict([('F1',
    0.9790794979079498),
    ('Precision', 0.975),
    ('Recall', 0.9831932773109243),
    ('AUC', 0.9981062847674281),
    ('Log Loss', 0.08064971198327839),
    ('MCC', 0.9436801731761278),
    ('ROC',
    (array([0.          , 0.          , 0.          , 0.          , 0.          ,

```

(continues on next page)

(continues on next page)

(continued from previous page)

```

0.71428571, 0.74789916, 0.76470588, 0.80672269, 0.82352941,
0.92436975, 0.92436975, 0.94117647, 0.94117647, 0.95798319,
0.95798319, 0.96638655, 0.96638655, 0.97478992, 0.97478992,
0.99159664, 0.99159664, 1., 1., 1.,
1., 1., 1., 1., 1.,
1., 1., 1., 1., 1.,
1. ]),
array([1.9917016 , 0.99170154, 0.99118173, 0.99113 , 0.99107367,
0.99095434, 0.9903882 , 0.9900671 , 0.9890872 , 0.9874872 ,
0.9870094 , 0.9863927 , 0.9848595 , 0.98429155, 0.9837299 ,
0.98244077, 0.97697645, 0.9723681 , 0.9694117 , 0.96865326,
0.96827 , 0.9661397 , 0.965724 , 0.96271837, 0.95904344,
0.95776254, 0.9541006 , 0.9514202 , 0.9272428 , 0.9244033 ,
0.729121 , 0.7108931 , 0.66336167, 0.6254638 , 0.5943486 ,
0.44532707, 0.372028 , 0.3703832 , 0.36668354, 0.28058025,
0.21105221, 0.1829909 , 0.18207934, 0.09410949, 0.07737678,
0.06360406, 0.06231202, 0.03264038, 0.02874083, 0.02029958,
0.01903956, 0.0164301 , 0.01615754, 0.01451924, 0.01421483,
0.01351323], dtype=float32))),
('Confusion Matrix', 0 1
0 65 6
1 5 114),
('# Training', 379),
('# Testing', 190)],
'score': 0.9539748953974896},
{'all_objective_scores': OrderedDict([('F1', 0.9451476793248945),
('Precision', 0.9491525423728814),
('Recall', 0.9411764705882353),
('AUC', 0.9807923169267707),
('Log Loss', 0.15895638581877822),
('MCC', 0.8530080688400891),
('ROC',
(array([0. , 0. , 0. , 0. , 0. ,
0. , 0. , 0. , 0. , 0. ,
0. , 0. , 0. , 0. , 0. ,
0.01428571, 0.01428571, 0.01428571, 0.01428571, 0.01428571,
0.01428571, 0.01428571, 0.01428571, 0.02857143, 0.02857143,
0.04285714, 0.04285714, 0.05714286, 0.05714286, 0.07142857,
0.07142857, 0.08571429, 0.08571429, 0.11428571, 0.11428571,
0.14285714, 0.15714286, 0.22857143, 0.25714286, 0.38571429,
0.44285714, 0.54285714, 0.6 , 0.64285714, 0.67142857,
0.7 , 1. ]),
array([0. , 0.18487395, 0.22689076, 0.23529412, 0.26890756,
0.29411765, 0.31092437, 0.33613445, 0.34453782, 0.36134454,
0.3697479 , 0.40336134, 0.42857143, 0.44537815, 0.46218487,
0.47058824, 0.50420168, 0.5210084 , 0.53781513, 0.55462185,
0.59663866, 0.65546218, 0.70588235, 0.70588235, 0.78991597,
0.78991597, 0.90756303, 0.90756303, 0.91596639, 0.91596639,
0.94117647, 0.94117647, 0.97478992, 0.97478992, 0.99159664,
0.99159664, 1. , 1. , 1. , 1. , 1. ,
1. , 1. , 1. , 1. , 1. ,
1. , 1. ]),
array([1.9912372 , 0.99123716, 0.99071807, 0.9905557 , 0.989953 ,
0.98967326, 0.9892195 , 0.9890625 , 0.9883869 , 0.9881627 ,
0.98800814, 0.9873002 , 0.9866671 , 0.9863214 , 0.98610663,
0.98588854, 0.9824102 , 0.9817692 , 0.9788537 , 0.9763956 ,
0.9736345 , 0.9734056 , 0.9697179 , 0.96919906, 0.9408594 ,

```

(continues on next page)

(continued from previous page)

```

0.93151474, 0.70544225, 0.6761148 , 0.6481446 , 0.56992525,
0.5188037 , 0.514126 , 0.40845907, 0.35379264, 0.33419853,
0.22962442, 0.22778042, 0.09829229, 0.05558122, 0.04127952,
0.03519488, 0.01854065, 0.01816005, 0.01665189, 0.01571225,
0.01527028, 0.01430225], dtype=float32))),
('Confusion Matrix',
 0 1
 0 64 6
 1 7 112),
('# Training', 380),
('# Testing', 189)]),
'score': 0.9451476793248945}],
1: {'id': 1,
'pipeline_class_name': 'XGBoostPipeline',
'pipeline_name': 'XGBoost Classifier w/ One Hot Encoder + Simple Imputer + RF_
Classifier Select From Model',
'parameters': {'eta': 0.38438170729269994,
'min_child_weight': 3.677811458900251,
'max_depth': 13,
'impute_strategy': 'median',
'percent_features': 0.793807787701838},
'score': 0.9648960642685166,
'high_variance_cv': False,
'training_time': 2.482739210128784,
'cv_data': [{'all_objective_scores': OrderedDict([('F1', 0.975),
('Precision', 0.9669421487603306),
('Recall', 0.9831932773109243),
('AUC', 0.9966859983429992),
('Log Loss', 0.07756886413615001),
('MCC', 0.932389423285531),
('ROC',
(array([0.
0.01408451, 0.01408451, 0.02816901, 0.02816901, 0.08450704,
0.08450704, 0.09859155, 0.09859155, 0.91549296, 0.94366197,
1.
])),
array([0.
0.00840336, 0.32773109, 0.34453782, 0.87394958,
0.87394958, 0.96638655, 0.96638655, 0.98319328, 0.98319328,
0.99159664, 0.99159664, 1.
, 1.
, 1.
,
1.
])),
array([1.9981308 , 0.99813086, 0.9959843 , 0.99595356, 0.9659731 ,
0.9646539 , 0.8474301 , 0.8226159 , 0.7340414 , 0.43467796,
0.35439238, 0.3156159 , 0.2735931 , 0.00334718, 0.00322822,
0.00207397], dtype=float32))],
('Confusion Matrix',
 0 1
 0 67 4
 1 2 117),
('# Training', 379),
('# Testing', 190)]),
'score': 0.975},
{'all_objective_scores': OrderedDict([('F1', 0.9661016949152542),
('Precision', 0.9743589743589743),
('Recall', 0.957983193277311),
('AUC', 0.9899396378269618),
('Log Loss', 0.1157939842209759),
('MCC', 0.9107843947529072),
('ROC',
(array([0.
0.
, 0.
, 0.
, 0.
, 0.
,
0.
, 0.
, 0.
, 0.
, 0.
, 0.01408451,

```

(continues on next page)

(continued from previous page)

```

0.01408451, 0.02816901, 0.02816901, 0.04225352, 0.04225352,
0.05633803, 0.05633803, 0.14084507, 0.14084507, 0.16901408,
0.16901408, 0.22535211, 0.49295775, 0.52112676, 0.66197183,
0.73239437, 0.81690141, 0.84507042, 0.87323944, 0.90140845,
0.97183099, 1.          ],
array([0.          , 0.01680672, 0.08403361, 0.11764706, 0.24369748,
0.2605042 , 0.26890756, 0.29411765, 0.57983193, 0.57983193,
0.91596639, 0.91596639, 0.94957983, 0.94957983, 0.95798319,
0.95798319, 0.98319328, 0.98319328, 0.99159664, 0.99159664,
1.          , 1.          , 1.          , 1.          , 1.          ,
1.          , 1.          , 1.          , 1.          , 1.          ,
1.          , 1.          ],
array([1.9985983e+00, 9.9859840e-01, 9.9832827e-01, 9.9818736e-01,
9.9679202e-01, 9.9674618e-01, 9.9667656e-01, 9.9658316e-01,
9.9108350e-01, 9.9083334e-01, 8.2265288e-01, 8.2121074e-01,
7.2431225e-01, 6.3007897e-01, 5.7422268e-01, 4.5153135e-01,
3.5701558e-01, 1.7140125e-01, 1.4191566e-01, 9.0782769e-02,
6.7663342e-02, 5.5308960e-02, 4.5769266e-03, 4.5044953e-03,
3.2702754e-03, 3.2080656e-03, 2.9759661e-03, 2.8877873e-03,
2.4864469e-03, 2.3314529e-03, 1.5917335e-03, 1.3716089e-
↪03],
dtype=float32))),
('Confusion Matrix',      0      1
0  68      3
1   5  114),
('# Training', 379),
('# Testing', 190)]),
'score': 0.9661016949152542},
{'all_objective_scores': OrderedDict([('F1', 0.9535864978902954),
('Precision', 0.9576271186440678),
('Recall', 0.9495798319327731),
('AUC', 0.9863145258103241),
('Log Loss', 0.14040204006858248),
('MCC', 0.8756320549488144),
('ROC',
(array([0.          , 0.          , 0.          , 0.          , 0.          ,
0.          , 0.          , 0.          , 0.          , 0.          ,
0.          , 0.          , 0.          , 0.          , 0.01428571,
0.01428571, 0.02857143, 0.02857143, 0.04285714, 0.04285714,
0.08571429, 0.08571429, 0.1          , 0.1          , 0.11428571,
0.11428571, 0.54285714, 0.57142857, 0.58571429, 0.61428571,
0.68571429, 0.71428571, 0.8          , 1.          ])),
array([0.          , 0.00840336, 0.04201681, 0.05882353, 0.1512605 ,
0.21008403, 0.23529412, 0.30252101, 0.31932773, 0.33613445,
0.35294118, 0.62184874, 0.65546218, 0.72268908, 0.72268908,
0.73109244, 0.73109244, 0.79831933, 0.79831933, 0.94957983,
0.94957983, 0.96638655, 0.96638655, 0.97478992, 0.97478992,
1.          , 1.          , 1.          , 1.          , 1.          ,
1.          , 1.          , 1.          , 1.          ])),
array([1.998668 , 0.99866796, 0.99859375, 0.99848515, 0.99842656,
0.9980634 , 0.9980153 , 0.9971877 , 0.9971783 , 0.9970214 ,
0.99701583, 0.9893309 , 0.98911965, 0.9835843 , 0.9835355 ,
0.98161846, 0.9813965 , 0.9743532 , 0.97359157, 0.62105346,
0.47956562, 0.46517578, 0.42012912, 0.36913908, 0.25974 ,
0.16882712, 0.00341962, 0.00341539, 0.00333552, 0.00333485,
0.00295924, 0.00287893, 0.00215919, 0.00207186],
↪dtype=float32))),

```

(continues on next page)

(continued from previous page)

```

        ('Confusion Matrix',
         0  1
         0 65  5
         1  6 113),
        ('# Training', 380),
        ('# Testing', 189)]),
        'score': 0.9535864978902954}}],
2: {'id': 2,
    'pipeline_class_name': 'RFClassificationPipeline',
    'pipeline_name': 'Random Forest Classifier w/ One Hot Encoder + Simple Imputer +
↳RF Classifier Select From Model',
    'parameters': {'n_estimators': 569,
                    'max_depth': 22,
                    'impute_strategy': 'most_frequent',
                    'percent_features': 0.8593661614465293},
    'score': 0.973476521410252,
    'high_variance_cv': False,
    'training_time': 9.765936613082886,
    'cv_data': [{'all_objective_scores': OrderedDict([('F1',
                                                         0.9831932773109243),
                                                         ('Precision', 0.9831932773109243),
                                                         ('Recall', 0.9831932773109243),
                                                         ('AUC', 0.9977512131613208),
                                                         ('Log Loss', 0.08333209715129118),
                                                         ('MCC', 0.9550242632264173),
                                                         ('ROC',
                  (array([0.          , 0.          , 0.          , 0.          , 0.          ,
                           0.          , 0.          , 0.          , 0.          , 0.          ,
                           0.          , 0.          , 0.          , 0.          , 0.          ,
                           0.          , 0.01408451, 0.01408451, 0.09859155, 0.09859155,
                           0.12676056, 0.12676056, 0.4084507 , 0.43661972, 0.47887324,
                           0.54929577, 0.5915493 , 0.6056338 , 0.64788732, 0.76056338,
                           1.          ]),
                  array([0.          , 0.26890756, 0.39495798, 0.42016807, 0.49579832,
                           0.52941176, 0.57142857, 0.59663866, 0.64705882, 0.65546218,
                           0.68907563, 0.69747899, 0.73109244, 0.8487395 , 0.89915966,
                           0.95798319, 0.95798319, 0.98319328, 0.98319328, 0.99159664,
                           0.99159664, 1.          , 1.          , 1.          , 1.          ,
                           1.          , 1.          , 1.          , 1.          , 1.          ,
                           1.          ]),
                  array([2.00000000e+00, 1.00000000e+00, 9.98242531e-01, 9.
↳96485062e-01,
                           9.94727592e-01, 9.92970123e-01, 9.91212654e-01, 9.
↳89455185e-01,
                           9.87697715e-01, 9.85940246e-01, 9.84182777e-01, 9.
↳82425308e-01,
                           9.77152900e-01, 9.08611599e-01, 8.89279438e-01, 7.
↳75043937e-01,
                           7.39894552e-01, 7.13532513e-01, 3.69068541e-01, 3.
↳53251318e-01,
                           2.86467487e-01, 2.26713533e-01, 4.92091388e-02, 4.
↳56942004e-02,
                           4.39367311e-02, 1.23022847e-02, 1.05448155e-02, 5.
↳27240773e-03,
                           3.51493849e-03, 1.75746924e-03, 0.00000000e+00]))),
        ('Confusion Matrix',
         0  1
         0 69  2
         1  2 117),

```

(continues on next page)

(continued from previous page)

```

        ('# Training', 379),
        ('# Testing', 190)]),
    'score': 0.9831932773109243},
    {'all_objective_scores': OrderedDict([('F1', 0.9789029535864979),
        ('Precision', 0.9830508474576272),
        ('Recall', 0.9747899159663865),
        ('AUC', 0.9923067818676766),
        ('Log Loss', 0.1143023991450681),
        ('MCC', 0.9439989100828842),
        ('ROC',
            (array([0.          , 0.          , 0.          , 0.          , 0.          ,
                    0.          , 0.          , 0.          , 0.          , 0.          ,
                    0.          , 0.01408451, 0.01408451, 0.01408451, 0.01408451,
                    0.01408451, 0.01408451, 0.01408451, 0.01408451, 0.01408451,
                    0.01408451, 0.01408451, 0.01408451, 0.01408451, 0.01408451,
                    0.01408451, 0.01408451, 0.02816901, 0.02816901, 0.04225352,
                    0.04225352, 0.11267606, 0.11267606, 0.22535211, 0.25352113,
                    0.30985915, 0.33802817, 0.3943662 , 0.42253521, 0.43661972,
                    0.49295775, 0.52112676, 0.6056338 , 0.73239437, 1.          ],
                    dtype=float64,
                    ndim=32)}),
        array([0.          , 0.19327731, 0.27731092, 0.33613445, 0.3697479 ,
                0.45378151, 0.47058824, 0.50420168, 0.52941176, 0.56302521,
                0.57142857, 0.57142857, 0.59663866, 0.61344538, 0.62184874,
                0.63865546, 0.64705882, 0.66386555, 0.68907563, 0.70588235,
                0.73109244, 0.76470588, 0.81512605, 0.83193277, 0.84033613,
                0.85714286, 0.94117647, 0.94117647, 0.98319328, 0.98319328,
                0.99159664, 0.99159664, 1.          , 1.          , 1.          ,
                1.          , 1.          , 1.          , 1.          , 1.          ,
                1.          , 1.          , 1.          , 1.          , 1.          ],
                dtype=float64,
                ndim=32)}),
        array([2.00000000e+00, 1.00000000e+00, 9.98242531e-01, 9.
        ↪96485062e-01,
        ↪9.94727592e-01, 9.91212654e-01, 9.89455185e-01, 9.
        ↪87697715e-01,
        ↪9.85940246e-01, 9.80667838e-01, 9.78910369e-01, 9.
        ↪77152900e-01,
        ↪9.75395431e-01, 9.73637961e-01, 9.71880492e-01, 9.
        ↪70123023e-01,
        ↪9.68365554e-01, 9.63093146e-01, 9.52548330e-01, 9.
        ↪50790861e-01,
        ↪9.36731107e-01, 9.29701230e-01, 8.82249561e-01, 8.
        ↪73462214e-01,
        ↪8.71704745e-01, 8.55887522e-01, 6.67838313e-01, 6.
        ↪55536028e-01,
        ↪4.95606327e-01, 4.81546573e-01, 4.69244288e-01, 3.
        ↪98945518e-01,
        ↪3.72583480e-01, 1.15992970e-01, 1.00175747e-01, 6.
        ↪50263620e-02,
        ↪4.39367311e-02, 2.81195079e-02, 1.75746924e-02, 1.
        ↪40597540e-02,
        ↪8.78734622e-03, 7.02987698e-03, 3.51493849e-03, 1.
        ↪75746924e-03,
            0.00000000e+00]))),
    ('Confusion Matrix',
     0  69  2
     1  3 116),
    ('# Training', 379),

```

(continues on next page)

(continued from previous page)

```

        ('# Testing', 190)]),
        'score': 0.9789029535864979},
    {'all_objective_scores': OrderedDict([('F1', 0.9583333333333334),
        ('Precision', 0.9504132231404959),
        ('Recall', 0.9663865546218487),
        ('AUC', 0.9847539015606243),
        ('Log Loss', 0.1416407874804379),
        ('MCC', 0.8861141678760591),
        ('ROC',
         (array([0.          , 0.          , 0.01428571, 0.01428571, 0.01428571,
                  0.01428571, 0.01428571, 0.01428571, 0.01428571, 0.01428571,
                  0.01428571, 0.02857143, 0.02857143, 0.02857143, 0.02857143,
                  0.04285714, 0.04285714, 0.08571429, 0.1          , 0.1          ,
                  0.14285714, 0.14285714, 0.37142857, 0.4          , 0.47142857,
                  0.5          , 0.51428571, 0.62857143, 1.          ]),
          array([0.          , 0.37815126, 0.46218487, 0.54621849, 0.57983193,
                  0.62184874, 0.6302521 , 0.65546218, 0.66386555, 0.69747899,
                  0.78151261, 0.78991597, 0.85714286, 0.87394958, 0.88235294,
                  0.88235294, 0.96638655, 0.96638655, 0.97478992, 0.99159664,
                  0.99159664, 1.          , 1.          , 1.          , 1.          ,
                  1.          , 1.          , 1.          , 1.          ]),
          array([2.00000000e+00, 1.00000000e+00, 9.98242531e-01, 9.
↪ 96485062e-01,
                  9.94727592e-01, 9.92970123e-01, 9.91212654e-01, 9.
↪ 89455185e-01,
                  9.87697715e-01, 9.85940246e-01, 9.43760984e-01, 9.
↪ 38488576e-01,
                  8.64674868e-01, 8.03163445e-01, 7.62741652e-01, 7.
↪ 34622144e-01,
                  5.55360281e-01, 5.04393673e-01, 4.60456942e-01, 3.
↪ 81370826e-01,
                  3.04042179e-01, 2.93497364e-01, 2.10896309e-02, 1.
↪ 93321617e-02,
                  8.78734622e-03, 5.27240773e-03, 3.51493849e-03, 1.
↪ 75746924e-03,
                  0.00000000e+00]))),
        ('Confusion Matrix',
         [[0 64 6
            1 4 115]]),
        ('# Training', 380),
        ('# Testing', 189)]),
        'score': 0.9583333333333334}}],
    3: {'id': 3,
        'pipeline_class_name': 'XGBoostPipeline',
        'pipeline_name': 'XGBoost Classifier w/ One Hot Encoder + Simple Imputer + RF_
↪ Classifier Select From Model',
        'parameters': {'eta': 0.5288949197529046,
                        'min_child_weight': 6.112401049845392,
                        'max_depth': 6,
                        'impute_strategy': 'most_frequent',
                        'percent_features': 0.34402219881309576},
        'score': 0.9485662145040741,
        'high_variance_cv': False,
        'training_time': 2.421769857406616,
        'cv_data': [{'all_objective_scores': OrderedDict([('F1',
                        0.9666666666666667),
                        ('Precision', 0.9586776859504132),

```

(continues on next page)

(continued from previous page)

```

('Recall', 0.9747899159663865),
('AUC', 0.9950289975144987),
('Log Loss', 0.09729690174817254),
('MCC', 0.9097672817424011),
('ROC',
 (array([0.          , 0.          , 0.          , 0.          , 0.          ,
         0.          , 0.          , 0.          , 0.          , 0.          ,
         0.          , 0.          , 0.          , 0.          , 0.          ,
         0.01408451, 0.01408451, 0.04225352, 0.04225352, 0.05633803,
         0.05633803, 0.09859155, 0.09859155, 0.12676056, 0.12676056,
         0.16901408, 0.16901408, 0.38028169, 0.4084507 , 0.47887324,
         0.52112676, 0.54929577, 0.57746479, 0.6056338 , 0.66197183,
         0.67605634, 0.98591549, 1.          ]),
  array([0.          , 0.02521008, 0.03361345, 0.16806723, 0.17647059,
         0.20168067, 0.21008403, 0.35294118, 0.36134454, 0.38655462,
         0.42016807, 0.43697479, 0.4789916 , 0.49579832, 0.51260504,
         0.54621849, 0.57142857, 0.57983193, 0.59663866, 0.91596639,
         0.91596639, 0.94957983, 0.94957983, 0.96638655, 0.96638655,
         0.97478992, 0.97478992, 0.98319328, 0.98319328, 0.99159664,
         0.99159664, 1.          , 1.          , 1.          , 1.          ,
         1.          , 1.          , 1.          , 1.          , 1.          ,
         1.          , 1.          , 1.          ]),
  array([1.9918091 , 0.99180907, 0.9914386 , 0.99143535, 0.9909027 ,
         0.9905533 , 0.9904917 , 0.9904881 , 0.990333 , 0.99021775,
         0.9898657 , 0.98977226, 0.9895096 , 0.9892888 , 0.98913723,
         0.98872095, 0.9886431 , 0.9884277 , 0.9880271 , 0.9180971 ,
         0.907636 , 0.7645348 , 0.63884944, 0.6283931 , 0.5994021 ,
         0.5585935 , 0.38392243, 0.37918517, 0.2239313 , 0.22208475,
         0.12249093, 0.12071476, 0.02558885, 0.0240444 , 0.01662812,
         0.01431511, 0.01352188, 0.01269417, 0.01247706, 0.01176526,
         0.01091322, 0.01024354, 0.0058776 ] , dtype=float32))),
('Confusion Matrix',
 0      1
0 66    5
1   3 116),
('# Training', 379),
('# Testing', 190))),
'score': 0.9666666666666667},
{'all_objective_scores': OrderedDict([('F1', 0.9338842975206612),
('Precision', 0.9186991869918699),
('Recall', 0.9495798319327731),
('AUC', 0.9846727423363711),
('Log Loss', 0.17591368435195795),
('MCC', 0.8188772553918354),
('ROC',
 (array([0.          , 0.          , 0.          , 0.          , 0.          ,
         0.          , 0.          , 0.          , 0.          , 0.          ,
         0.          , 0.          , 0.01408451, 0.01408451, 0.01408451,
         0.01408451, 0.01408451, 0.02816901, 0.02816901, 0.04225352, 0.07042254, 0.07042254,
         0.08450704, 0.08450704, 0.15492958, 0.15492958, 0.18309859, 0.18309859,
         0.22535211, 0.22535211, 0.35211268, 0.38028169, 0.4084507 ,
         0.46478873, 0.49295775, 0.50704225, 0.78873239, 0.8028169 ,
         0.84507042, 1.          ]),
  array([0.          , 0.10084034, 0.11764706, 0.14285714, 0.19327731,
         0.22689076, 0.4789916 , 0.49579832, 0.53781513, 0.58823529,
         0.60504202, 0.66386555, 0.68067227, 0.69747899, 0.73109244,
```

(continues on next page)

(continued from previous page)

```

0.75630252, 0.78991597, 0.80672269, 0.87394958, 0.87394958,
0.91596639, 0.92436975, 0.92436975, 0.94117647, 0.94117647,
0.94957983, 0.94957983, 0.96638655, 0.96638655, 0.99159664,
0.99159664, 1., 1., 1., 1., 1.,
1., 1., 1., 1., 1.,
1., 1. ]),
array([1.9947486 , 0.99474853, 0.99465054, 0.9945134 , 0.9920238 ,
0.99128073, 0.9897943 , 0.98960495, 0.98933965, 0.98638594,
0.98634803, 0.9774489 , 0.9764572 , 0.97513086, 0.9740399 ,
0.9718917 , 0.9563253 , 0.946435 , 0.8835335 , 0.878611 ,
0.7919501 , 0.755015 , 0.71543545, 0.70243114, 0.6865198 ,
0.64198273, 0.47787225, 0.22530945, 0.16627482, 0.08063855,
0.06376061, 0.04771832, 0.01834523, 0.01635974, 0.01409694,
0.01397619, 0.01379675, 0.01324805, 0.01214174, 0.00822058,
0.0067653 , 0.00625362], dtype=float32))),
('Confusion Matrix', 0 1
0 61 10
1 6 113),
('# Training', 379),
('# Testing', 190)]),
'score': 0.9338842975206612},
{'all_objective_scores': OrderedDict([('F1', 0.9451476793248945),
('Precision', 0.9491525423728814),
('Recall', 0.9411764705882353),
('AUC', 0.9645258103241297),
('Log Loss', 0.18750503725032247),
('MCC', 0.8530080688400891),
('ROC',
(array([0. , 0. , 0. , 0.01428571, 0.02857143,
0.02857143, 0.02857143, 0.02857143, 0.02857143, 0.02857143,
0.02857143, 0.02857143, 0.02857143, 0.02857143, 0.02857143,
0.04285714, 0.04285714, 0.05714286, 0.05714286, 0.07142857,
0.07142857, 0.08571429, 0.08571429, 0.1 , 0.1 ,
0.12857143, 0.12857143, 0.17142857, 0.17142857, 0.18571429,
0.18571429, 0.41428571, 0.44285714, 0.47142857, 0.5 ,
0.64285714, 1. ]),
array([0. , 0.00840336, 0.01680672, 0.01680672, 0.46218487,
0.48739496, 0.5210084 , 0.54621849, 0.55462185, 0.58823529,
0.61344538, 0.6302521 , 0.63865546, 0.65546218, 0.77310924,
0.77310924, 0.8487395 , 0.8487395 , 0.87394958, 0.87394958,
0.93277311, 0.93277311, 0.94957983, 0.94957983, 0.96638655,
0.97478992, 0.98319328, 0.98319328, 0.99159664, 0.99159664,
1. , 1. , 1. , 1. , 1. ,
1. , 1. ]),
array([1.9947662 , 0.99476624, 0.99434704, 0.99431866, 0.993646 ,
0.99334115, 0.9933149 , 0.99291396, 0.99289966, 0.9928403 ,
0.9924971 , 0.9922282 , 0.99102247, 0.9907026 , 0.9769703 ,
0.9766011 , 0.9003137 , 0.89544696, 0.88858014, 0.8658668 ,
0.68705285, 0.6113713 , 0.43407622, 0.37409404, 0.32079345,
0.29611424, 0.29083598, 0.2146415 , 0.21115056, 0.1591847 ,
0.14415835, 0.01072225, 0.01065906, 0.00977037, 0.00951847,
0.00907114, 0.00867252], dtype=float32))),
('Confusion Matrix', 0 1
0 64 6
1 7 112),
('# Training', 380),
('# Testing', 189)]),

```

(continues on next page)

(continued from previous page)

```

        'score': 0.9451476793248945}}],
4: {'id': 4,
    'pipeline_class_name': 'LogisticRegressionPipeline',
    'pipeline_name': 'Logistic Regression Classifier w/ One Hot Encoder + Simple_
↳Imputer + Standard Scaler',
    'parameters': {'penalty': 'l2',
                    'C': 8.444214828324364,
                    'impute_strategy': 'most_frequent'},
    'score': 0.9719486068554629,
    'high_variance_cv': False,
    'training_time': 1.6746852397918701,
    'cv_data': [{'all_objective_scores': OrderedDict([('F1',
0.9709543568464729),
('Precision', 0.9590163934426229),
('Recall', 0.9831932773109243),
('AUC', 0.9875724937862469),
('Log Loss', 0.14613712409589413),
('MCC', 0.9211492315750531),
('ROC',
(array([0.          , 0.          , 0.          , 0.01408451, 0.01408451,
0.02816901, 0.02816901, 0.04225352, 0.04225352, 0.05633803,
0.05633803, 0.07042254, 0.07042254, 0.08450704, 0.08450704,
1.          ]),
array([0.          , 0.00840336, 0.57983193, 0.57983193, 0.68067227,
0.68067227, 0.92436975, 0.92436975, 0.95798319, 0.95798319,
0.98319328, 0.98319328, 0.99159664, 0.99159664, 1.          ,
1.          ]),
array([2.00000000e+00, 1.00000000e+00, 9.99936259e-01, 9.
↳99935430e-01,
          9.99856221e-01, 9.99847656e-01, 9.44959587e-01, 9.
↳25157910e-01,
          8.68396485e-01, 8.05626231e-01, 6.85798432e-01, 5.
↳24509038e-01,
          4.79378624e-01, 2.63698337e-01, 2.31528899e-01, 1.
↳13164746e-48]))),
    ('Confusion Matrix',
     0    1
     0 66   5
     1  2 117),
    ('# Training', 379),
    ('# Testing', 190)]),
    'score': 0.9709543568464729},
{'all_objective_scores': OrderedDict([('F1', 0.9790794979079498),
('Precision', 0.975),
('Recall', 0.9831932773109243),
('AUC', 0.9910048526452835),
('Log Loss', 0.12072855977761998),
('MCC', 0.9436801731761278),
('ROC',
(array([0.          , 0.          , 0.          , 0.01408451, 0.01408451,
0.05633803, 0.05633803, 0.23943662, 0.23943662, 1.          ,
↳]),
array([0.          , 0.00840336, 0.5210084 , 0.5210084 , 0.98319328,
0.98319328, 0.99159664, 0.99159664, 1.          , 1.          ,
↳]),
array([2.00000000e+00, 1.00000000e+00, 9.99838725e-01, 9.
↳99835686e-01,
          6.74763310e-01, 4.29302726e-01, 3.81937028e-01, 5.
↳59186779e-04,

```

(continues on next page)

(continued from previous page)

```

                2.40539384e-04, 1.23958997e-29]])),
('Confusion Matrix',
 0      1
 0  68   3
 1   2 117),
('# Training', 379),
('# Testing', 190)],
'score': 0.9790794979079498},
{'all_objective_scores': OrderedDict([('F1', 0.9658119658119659),
('Precision', 0.9826086956521739),
('Recall', 0.9495798319327731),
('AUC', 0.9965186074429772),
('Log Loss', 0.08056331735168969),
('MCC', 0.9112159507396058),
('ROC',
(array([0.          , 0.          , 0.          , 0.01428571, 0.01428571,
0.02857143, 0.02857143, 0.04285714, 0.04285714, 0.05714286,
0.05714286, 0.08571429, 0.08571429, 1.          ])),
array([0.          , 0.00840336, 0.90756303, 0.90756303, 0.93277311,
0.93277311, 0.94957983, 0.94957983, 0.98319328, 0.98319328,
0.99159664, 0.99159664, 1.          , 1.          ])),
array([2.00000000e+00, 9.99999999e-01, 9.46260690e-01, 9.
↪45033115e-01,
          9.06023542e-01, 8.77670109e-01, 7.39987181e-01, 4.
↪77332145e-01,
          3.62040644e-01, 2.35727064e-01, 1.93164591e-01, 9.
↪53728712e-02,
          8.79204445e-02, 1.15873760e-20]])),
('Confusion Matrix',
 0      1
 0  68   2
 1   6 113),
('# Training', 380),
('# Testing', 189)],
'score': 0.9658119658119659}}],
'search_order': [0, 1, 2, 3, 4]}

```

2.6.7 Regression Example

```

[1]: import evalml
from evalml import AutoRegressionSearch
from evalml.demos import load_diabetes
from evalml.pipelines import PipelineBase, get_pipelines

X, y = evalml.demos.load_diabetes()

automl = AutoRegressionSearch(objective="R2", max_pipelines=5)

automl.search(X, y)

*****
* Beginning pipeline search *
*****

Optimizing for R2. Greater score is better.

```

(continues on next page)

(continued from previous page)

```
Searching up to 5 pipelines.
Possible model types: linear_model, random_forest
```

```
FigureWidget({
  'data': [{ 'mode': 'lines+markers',
             'name': 'Best Score',
             'type'...
```

```
Random Forest Regressor w/ One Hot ... 20%|          | Elapsed:00:09
Random Forest Regressor w/ One Hot ... 40%|          | Elapsed:00:16
Linear Regressor w/ One Hot Encoder... 60%|          | Elapsed:00:16
Random Forest Regressor w/ One Hot ... 80%|          | Elapsed:00:24
Random Forest Regressor w/ One Hot ... 100%|         | Elapsed:00:35
Optimization finished                  100%|         | Elapsed:00:35
```

```
[2]: automl.rankings
```

```
[2]:   id      pipeline_class_name      score  high_variance_cv  \
0    2  LinearRegressionPipeline  0.488703             False
1    0    RFRegressionPipeline   0.422322             False
2    4    RFRegressionPipeline   0.391463             False
3    3    RFRegressionPipeline   0.383134             False
4    1    RFRegressionPipeline   0.381204             False
```

```
                                parameters
0  {'impute_strategy': 'mean', 'normalize': True,...
1  {'n_estimators': 569, 'max_depth': 22, 'impute...
2  {'n_estimators': 715, 'max_depth': 7, 'impute...
3  {'n_estimators': 609, 'max_depth': 7, 'impute...
4  {'n_estimators': 369, 'max_depth': 10, 'impute...
```

```
[3]: automl.best_pipeline
```

```
[3]: <evalml.pipelines.regression.linear_regression.LinearRegressionPipeline at 0x7fa7ac8eceb8>
```

```
[4]: automl.get_pipeline(0)
```

```
[4]: <evalml.pipelines.regression.random_forest.RFRegressionPipeline at 0x7fa7ad95df28>
```

```
[5]: automl.describe_pipeline(0)
```

```
*****
* Random Forest Regressor w/ One Hot Encoder + Simple Imputer + RF Regressor Select_
↳ From Model *
*****

Problem Types: Regression
Model Type: Random Forest
Objective to Optimize: R2 (greater is better)
Number of features: 8

Pipeline Steps
=====
1. One Hot Encoder
2. Simple Imputer
   * impute_strategy : most_frequent
```

(continues on next page)

(continued from previous page)

```

3. RF Regressor Select From Model
    * percent_features : 0.8593661614465293
    * threshold : -inf
4. Random Forest Regressor
    * n_estimators : 569
    * max_depth : 22

Training
=====
Training for Regression problems.
Total training time (including CV): 9.8 seconds

Cross Validation
-----

```

	R2	MAE	MSE	MedianAE	MaxError	ExpVariance	# Training	#_
↪Testing								
0	0.427	46.033	3276.018	39.699	161.858	0.428	294.000	148.
↪000								
1	0.450	48.953	3487.566	44.344	160.513	0.451	295.000	147.
↪000								
2	0.390	47.401	3477.117	41.297	171.420	0.390	295.000	147.
↪000								
mean	0.422	47.462	3413.567	41.780	164.597	0.423	-	↪
↪-								
std	0.031	1.461	119.235	2.360	5.947	0.031	-	↪
↪-								
coef of var	0.072	0.031	0.035	0.056	0.036	0.073	-	↪
↪-								

2.6.8 EvalML Components and Pipelines

EvalML searches and trains multiple machine learning **pipelines** in order to find the best one for your data. Each pipeline is made up of various **components** that can learn from the data, transform the data and ultimately predict labels given new data. Below we'll show an example of an EvalML pipeline. You can find a more in-depth look into [components](#) or learn how you can construct and use your own [pipelines](#).

XGBoost Pipeline

The EvalML XGBoost Pipeline is made up of four different components: a one-hot encoder, a missing value imputer, a feature selector and an XGBoost estimator. We can see them here by calling `.plot()`:

```

[1]: from evalml.pipelines import XGBoostPipeline

xgp = XGBoostPipeline(objective='recall', eta=0.5, min_child_weight=5, max_depth=10,
↪impute_strategy='mean', percent_features=0.5, number_features=10)
xgp.plot()

```

[1]: From the above graph we can see each component and its parameters. Each component takes in data and feeds it to the next. You can see more detailed information by calling `.describe()`:

```

[2]: xgp.describe()

```

```
*****
* XGBoost Classifier w/ One Hot Encoder + Simple Imputer + RF Classifier Select From_
  ↳Model *
*****

Problem Types: Binary Classification, Multiclass Classification
Model Type: XGBoost Classifier
Objective to Optimize: Recall (greater is better)

Pipeline Steps
=====
1. One Hot Encoder
2. Simple Imputer
   * impute_strategy : mean
3. RF Classifier Select From Model
   * percent_features : 0.5
   * threshold : -inf
4. XGBoost Classifier
   * eta : 0.5
   * max_depth : 10
   * min_child_weight : 5
```

You can then fit and score an individual pipeline:

```
[3]: import evalml

X, y = evalml.demos.load_breast_cancer()
xgp.fit(X, y)

xgp.score(X, y)

[3]: (0.9719887955182073, {})
```

2.6.9 EvalML Components

From the [overview](#), we see how each machine learning pipeline consists of individual components that process data before the data is ultimately sent to an estimator. Below we will describe each type of component in an EvalML pipeline.

Component Classes

Components can be split into two distinct classes: **transformers** and **estimators**.

```
[1]: import numpy as np
import pandas as pd
from evalml.pipelines.components import SimpleImputer

X = pd.DataFrame([[1, 2, 3], [1, np.nan, 3]])
display(X)

   0    1    2
0  1  2.0  3
1  1  NaN  3
```

Transformers take in data as input and output altered data. For example, an *imputer* takes in data and outputs filled in missing data with the mean, median, or most frequent value of each column.

A transformer can fit on data and then transform it in two steps by calling `.fit()` and `.transform()` or in one step by calling `fit_transform()`.

```
[2]: imp = SimpleImputer(impute_strategy="mean")
X = imp.fit_transform(X)

display(X)

   0   1   2
0  1  2.0  3
1  1  2.0  3
```

On the other hand, an estimator fits on data (X) and labels (y) in order to take in new data as input and return the predicted label as output. Therefore, an estimator can fit on data and labels by calling `.fit()` and then predict by calling `.predict()` on new data. An example of this would be the [*LogisticRegressionClassifier*](#). We can now see how a transformer alters data to make it easier for an estimator to learn and predict.

```
[3]: from evalml.pipelines.components import LogisticRegressionClassifier

clf = LogisticRegressionClassifier()

X = X
y = [1, 0]

clf.fit(X, y)
clf.predict(X)

[3]: array([0, 0])
```

Component Types

Components can further separate into different types that serve different functionality. Below we will go over the different types of transformers and estimators.

Transformer Types

- Imputer: fills missing data
 - Ex: [*SimpleImputer*](#)
- Scaler: alters numerical data into different scales
 - Ex: [*StandardScaler*](#)
- Encoder: translates different data types
 - Ex: [*OneHotEncoder*](#)
- Feature Selection: selects most useful columns of data
 - Ex: [*RFClassifierSelectFromModel*](#)

Estimator Types

- Regressor: predicts numerical or continuous labels
 - Ex: [*LinearRegressor*](#)

- Classifier: predicts categorical or discrete labels
 - Ex: *XGBoostClassifier*

2.6.10 Custom Pipelines in EvalML

EvalML pipelines consist of modular components combining any number of transformers and an estimator. This allows you to create pipelines that fit the needs of your data to achieve the best results. You can create your own pipeline like this:

```
[1]: from evalml.pipelines import PipelineBase
      from evalml.pipelines.components import StandardScaler

      # objectives can be either a str or the evalml objective object
      objective = 'Precision_Macro'

      # components can be passed in as a ComponentType/str, Component Object, or str of a
      ↪Component name
      # the components below show the different ways components can be passed in
      component_list = ['imputer', StandardScaler(), 'Logistic Regression Classifier']
      pipeline = PipelineBase(objective, component_list, n_jobs=-1, random_state=0)
```

```
[2]: from evalml.demos import load_wine

      X, y = load_wine()

      pipeline.fit(X, y)
      pipeline.score(X, y)

[2]: (1.0, {})
```

2.6.11 Guardrails

EvalML provides guardrails to help guide you in achieving the highest performing model. These utility functions help deal with overfitting, abnormal data, and missing data. These guardrails can be found under `evalml/guardrails/` `utils`. Below we will cover abnormal and missing data guardrails. You can find an in-depth look into overfitting guardrails [here](#).

Missing Data

Missing data or rows with NaN values provide many challenges for machine learning pipelines. In the worst case, many algorithms simply will not run with missing data! EvalML pipelines contain imputation *components* to ensure that doesn't happen. Imputation works by approximating missing values with existing values. However, if a column contains a high number of missing values a large percentage of the column would be approximated by a small percentage. This could potentially create a column without useful information for machine learning pipelines. By running the `detect_highly_null()` guardrail, EvalML will alert you to this potential problem by returning the columns that pass the missing values threshold.

```
[1]: import numpy as np
      import pandas as pd

      from evalml.guardrails.utils import detect_highly_null
```

(continues on next page)

(continued from previous page)

```
X = pd.DataFrame(
    [
        [1, 2, 3],
        [0, 4, np.nan],
        [1, 4, np.nan],
        [9, 4, np.nan],
        [8, 6, np.nan]
    ]
)

detect_highly_null(X, percent_threshold=0.8)
```

```
[1]: {2: 0.8}
```

Abnormal Data

EvalML provides two utility functions to check for abnormal data: `detect_outliers()` and `detect_id_columns()`.

ID Columns

ID columns in your dataset provide little to no benefit to a machine learning pipeline as the pipeline cannot extrapolate useful information from unique identifiers. Thus, `detect_id_columns()` reminds you if these columns exists.

```
[2]: from evalml.guardrails.utils import detect_id_columns

X = pd.DataFrame([[0, 53, 6325, 5], [1, 90, 6325, 10], [2, 90, 18, 20]], columns=['user_
↪number', 'cost', 'revenue', 'id'])

display(X)
print(detect_id_columns(X, threshold=0.95))
```

	user_number	cost	revenue	id
0	0	53	6325	5
1	1	90	6325	10
2	2	90	18	20

```
{'id': 1.0, 'user_number': 0.95}
```

Outliers

Outliers are observations that differ significantly from other observations in the same sample. Many machine learning pipelines suffer in performance if outliers are not dropped from the training set as they are not representative of the data. `detect_outliers()` uses Isolation Forests to notify you if a sample can be considered an outlier.

Below we generate a random dataset with some outliers.

```
[3]: data = np.random.randn(100, 100)
X = pd.DataFrame(data=data)

# outliers
X.iloc[3, :] = pd.Series(np.random.randn(100) * 10)
```

(continues on next page)

(continued from previous page)

```
X.iloc[25, :] = pd.Series(np.random.randn(100) * 20)
X.iloc[55, :] = pd.Series(np.random.randn(100) * 100)
X.iloc[72, :] = pd.Series(np.random.randn(100) * 100)
```

We then utilize `detect_outliers` to rediscover these outliers.

```
[4]: from evalml.guardrails.utils import detect_outliers

detect_outliers(X)

[4]: [3, 25, 55, 72]
```

2.6.12 Avoiding Overfitting

The ultimate goal of machine learning is to make accurate predictions on unseen data. EvalML aims to help you build a model that will perform as you expect once it is deployed in to the real world.

One of the benefits of using EvalML to build models is that it provides guardrails to ensure you are building pipelines that will perform reliably in the future. This page describes the various ways EvalML helps you avoid overfitting to your data.

```
[1]: import evalml
```

Detecting Label Leakage

A common problem is having features that include information from your label in your training data. By default, EvalML will provide a warning when it detects this may be the case.

Let's set up a simple example to demonstrate what this looks like

```
[2]: import pandas as pd

X = pd.DataFrame({
    "leaked_feature": [6, 6, 10, 5, 5, 11, 5, 10, 11, 4],
    "leaked_feature_2": [3, 2.5, 5, 2.5, 3, 5.5, 2, 5, 5.5, 2],
    "valid_feature": [3, 1, 3, 2, 4, 6, 1, 3, 3, 11]
})

y = pd.Series([1, 1, 0, 1, 1, 0, 1, 0, 0, 1])

automl = evalml.AutoClassificationSearch(
    max_pipelines=1,
    model_types=["linear_model"],
)

automl.search(X, y)

*****
* Beginning pipeline search *
*****

Optimizing for Precision. Greater score is better.

Searching up to 1 pipelines.
Possible model types: linear_model
```

(continues on next page)

(continued from previous page)

```

WARNING: Possible label leakage: leaked_feature, leaked_feature_2

FigureWidget({
  'data': [{ 'mode': 'lines+markers',
             'name': 'Best Score',
             'type'...

Logistic Regression Classifier w/ O... 100%| Elapsed:00:01
Optimization finished                  100%| Elapsed:00:01

```

In the example above, EvalML warned about the input features `leaked_feature` and `leak_feature_2`, which are both very closely correlated with the label we are trying to predict. If you'd like to turn this check off, set `detect_label_leakage=False`.

The second way to find features that may be leaking label information is to look at the top features of the model. As we can see below, the top features in our model are the 2 leaked features.

```

[3]: best_pipeline = automl.best_pipeline
     best_pipeline.feature_importances

[3]:
      feature  importance
0  leaked_feature -1.721463
1  leaked_feature_2 -1.643452
2   valid_feature -0.541728

```

Perform cross-validation for pipeline evaluation

By default, EvalML performs 3-fold cross validation when building pipelines. This means that it evaluates each pipeline 3 times using different sets of data for training and testing. In each trial, the data used for testing has no overlap from the data used for training.

While this is a good baseline approach, you can pass your own cross validation object to be used during modeling. The cross validation object can be any of the CV methods defined in [scikit-learn](#) or use a compatible API.

For example, if we wanted to do a time series split:

```

[4]: from sklearn.model_selection import TimeSeriesSplit

X, y = evalml.demos.load_breast_cancer()

automl = evalml.AutoClassificationSearch(
    cv=TimeSeriesSplit(n_splits=6),
    max_pipelines=1
)

automl.search(X, y)

*****
* Beginning pipeline search *
*****

Optimizing for Precision. Greater score is better.

Searching up to 1 pipelines.
Possible model types: linear_model, xgboost, random_forest

```

```
FigureWidget({
  'data': [{ 'mode': 'lines+markers',
             'name': 'Best Score',
             'type'...

XGBoost Classifier w/ One Hot Encod... 100%| Elapsed:00:04
Optimization finished                  100%| Elapsed:00:04
```

if we describe the 1 pipeline we built, we can see the scores for each of the 6 splits as determined by the cross-validation object we provided. We can also see the number of training examples per fold increased because we were using `TimeSeriesSplit`

```
[5]: automl.describe_pipeline(0)

*****
* XGBoost Classifier w/ One Hot Encoder + Simple Imputer + RF Classifier Select From_
↪Model *
*****

Problem Types: Binary Classification, Multiclass Classification
Model Type: XGBoost Classifier
Objective to Optimize: Precision (greater is better)
Number of features: 18

Pipeline Steps
=====
1. One Hot Encoder
2. Simple Imputer
   * impute_strategy : most_frequent
3. RF Classifier Select From Model
   * percent_features : 0.6273280598181127
   * threshold : -inf
4. XGBoost Classifier
   * eta : 0.5928446182250184
   * max_depth : 4
   * min_child_weight : 8.598391737229157

Training
=====
Training for Binary Classification problems.
Total training time (including CV): 4.8 seconds

Cross Validation
-----
```

	Precision	F1	Recall	AUC	Log Loss	MCC	# Training	# Testing
0	0.974	0.822	0.712	0.950	0.578	0.650	83.000	81.000
1	1.000	0.988	0.977	1.000	0.163	0.976	164.000	81.000
2	0.981	0.981	0.981	0.968	0.139	0.944	245.000	81.000
3	0.963	0.929	0.897	0.991	0.113	0.774	326.000	81.000
4	0.984	0.960	0.938	0.993	0.147	0.830	407.000	81.000
5	0.983	0.983	0.983	0.998	0.083	0.936	488.000	81.000
mean	0.981	0.944	0.915	0.983	0.204	0.852	-	-
std	0.012	0.064	0.105	0.020	0.186	0.125	-	-
coef of var	0.013	0.067	0.115	0.020	0.909	0.147	-	-

Detect unstable pipelines

When we perform cross validation we are trying generate an estimate of pipeline performance. EvalML does this by taking the mean of the score across the folds. If the performance across the folds varies greatly, it is indicative the the estimated value may be unreliable.

To protect the user against this, EvalML checks to see if the pipeline’s performance has a variance between the different folds. EvalML triggers a warning if the “coefficient of variance” of the scores (the standard deviation divided by mean) of the pipelines scores exceeds .2.

This warning will appear in the pipeline rankings under `high_variance_cv`.

```
[6]: automl.rankings
[6]:   id pipeline_class_name      score  high_variance_cv  \
0    0      XGBoostPipeline  0.980845                False
                                     parameters
0  {'eta': 0.5928446182250184, 'min_child_weight'...
```

Create holdout for model validation

EvalML offers a method to quickly create an holdout validation set. A holdout validation set is data that is not used during the process of optimizing or training the model. You should only use this validation set once you’ve picked the final model you’d like to use.

Below we create a holdout set of 20% of our data

```
[7]: X, y = evalml.demos.load_breast_cancer()
X_train, X_holdout, y_train, y_holdout = evalml.preprocessing.split_data(X, y, test_
↪size=.2)
```

```
[8]: automl = evalml.AutoClassificationSearch(
      objective="recall",
      max_pipelines=3,
      detect_label_leakage=True
    )
automl.search(X_train, y_train)

*****
* Beginning pipeline search *
*****

Optimizing for Recall. Greater score is better.

Searching up to 3 pipelines.
Possible model types: linear_model, xgboost, random_forest

FigureWidget({
  'data': [{'mode': 'lines+markers',
            'name': 'Best Score',
            'type'...

XGBoost Classifier w/ One Hot Encod... 33%|      | Elapsed:00:02
XGBoost Classifier w/ One Hot Encod... 67%|     | Elapsed:00:04
Random Forest Classifier w/ One Hot... 100%|| Elapsed:00:14
Optimization finished                  100%|| Elapsed:00:14
```

then we can retrain the best pipeline on all of our training data and see how it performs compared to the estimate

```
[9]: pipeline = automl.best_pipeline
      pipeline.fit(X_train, y_train)
      pipeline.score(X_holdout, y_holdout)

[9]: (0.9444444444444444, {})
```

2.6.13 Changelog

Future Releases

- Enhancements
- Fixes
- Changes
- Documentation Changes
- Testing Changes

v0.6.0 Dec. 16, 2019

- **Enhancements**
 - Added ability to create a plot of feature importances [#133](#)
 - Add early stopping to AutoML using patience and tolerance parameters [#241](#)
 - Added ROC and confusion matrix metrics and plot for classification problems and introduce PipelineSearchPlots class [#242](#)
 - Enhanced AutoML results with search order [#260](#)
- **Fixes**
 - Lower botocore requirement [#235](#)
 - Fixed decision_function calculation for FraudCost objective [#254](#)
 - Fixed return value of Recall metrics [#264](#)
- **Changes**
 - Renamed automl classes to AutoRegressionSearch and AutoClassificationSearch [#287](#)
 - Updating demo datasets to retain column names [#223](#)
 - Moving pipeline visualization to PipelinePlots class [#228](#)
 - Standarizing inputs as pd.DataFrame / pd.Series [#130](#)
 - Enforcing that pipelines must have an estimator as last component [#277](#)
 - Added ipywidgets as a dependency in requirements.txt [#278](#)
- **Documentation Changes**
 - Adding class properties to API reference [#244](#)
 - Fix and filter FutureWarnings from scikit-learn [#249](#), [#257](#)
 - Adding Linear Regression to API reference and cleaning up some Sphinx warnings [#227](#)
- **Testing Changes**
 - Added support for testing on Windows with CircleCI [#226](#)

- Added support for doctests [#233](#)

Warning: Breaking Changes

- The `fit()` method for `AutoClassifier` and `AutoRegressor` has been renamed to `search()`.
- `AutoClassifier` has been renamed to `AutoClassificationSearch`
- `AutoRegressor` has been renamed to `AutoRegressionSearch`
- `AutoClassificationSearch.results` and `AutoRegressionSearch.results` now is a dictionary with `pipeline_results` and `search_order` keys. `pipeline_results` can be used to access a dictionary that is identical to the old `.results` dictionary. Whereas, `search_order` returns a list of the search order in terms of pipeline id.
- Pipelines now require an estimator as the last component in `component_list`. Slicing pipelines now throws an `NotImplementedError` to avoid returning Pipelines without an estimator.

v0.5.2 Nov. 18, 2019

- **Enhancements**
 - Adding basic pipeline structure visualization [#211](#)
- **Documentation Changes**
 - Added notebooks to build process [#212](#)

v0.5.1 Nov. 15, 2019

- **Enhancements**
 - Added basic outlier detection guardrail [#151](#)
 - Added basic ID column guardrail [#135](#)
 - Added support for unlimited pipelines with a `max_time` limit [#70](#)
 - Updated `.readthedocs.yaml` to successfully build [#188](#)
- **Fixes**
 - Removed MSLE from default additional objectives [#203](#)
 - Fixed `random_state` passed in pipelines [#204](#)
 - Fixed slow down in `RFRegressor` [#206](#)
- **Changes**
 - Pulled information for `describe_pipeline` from pipeline's new `describe` method [#190](#)
 - Refactored pipelines [#108](#)
 - Removed guardrails from `Auto(*)` [#202](#), [#208](#)
- **Documentation Changes**
 - Updated documentation to show `max_time` enhancements [#189](#)
 - Updated release instructions for RTD [#193](#)
 - Added notebooks to build process [#212](#)
 - Added contributing instructions [#213](#)
 - Added new content [#222](#)

v0.5.0 Oct. 29, 2019

- **Enhancements**

- Added basic one hot encoding [#73](#)
- Use enums for `model_type` [#110](#)
- Support for splitting regression datasets [#112](#)
- Auto-infer multiclass classification [#99](#)
- Added support for other units in `max_time` [#125](#)
- Detect highly null columns [#121](#)
- Added additional regression objectives [#100](#)
- Show an interactive iteration vs. score plot when using `fit()` [#134](#)

- **Fixes**

- Reordered `describe_pipeline` [#94](#)
- Added type check for `model_type` [#109](#)
- Fixed `s` units when setting string `max_time` [#132](#)
- Fix objectives not appearing in API documentation [#150](#)

- **Changes**

- Reorganized tests [#93](#)
- Moved logging to its own module [#119](#)
- Show progress bar history [#111](#)
- Using cloudpickle instead of pickle to allow unloading of custom objectives [#113](#)
- Removed `render.py` [#154](#)

- **Documentation Changes**

- Update release instructions [#140](#)
- Include `additional_objectives` parameter [#124](#)
- Added Changelog [#136](#)

- **Testing Changes**

- Code coverage [#90](#)
- Added CircleCI tests for other Python versions [#104](#)
- Added doc notebooks as tests [#139](#)
- Test metadata for CircleCI and 2 core parallelism [#137](#)

v0.4.1 Sep. 16, 2019

- **Enhancements**

- Added AutoML for classification and regressor using Autobase and Skopt [#7](#) [#9](#)
- Implemented standard classification and regression metrics [#7](#)
- Added logistic regression, random forest, and XGBoost pipelines [#7](#)
- Implemented support for custom objectives [#15](#)

- Feature importance for pipelines #18
- Serialization for pipelines #19
- Allow fitting on objectives for optimal threshold #27
- Added detect label leakage #31
- Implemented callbacks #42
- Allow for multiclass classification #21
- Added support for additional objectives #79
- **Fixes**
 - Fixed feature selection in pipelines #13
 - Made random_seed usage consistent #45
- **Documentation Changes**
 - Documentation Changes
 - Added docstrings #6
 - Created notebooks for docs #6
 - Initialized readthedocs EvalML #6
 - Added favicon #38
- **Testing Changes**
 - Added testing for loading data #39

v0.2.0 Aug. 13, 2019

- **Enhancements**
 - Created fraud detection objective #4

v0.1.0 July. 31, 2019

- *First Release*
- **Enhancements**
 - Added lead scoring objective #1
 - Added basic classifier #1
- **Documentation Changes**
 - Initialized Sphinx for docs #1

2.6.14 API Reference

Demo Datasets

<code>load_fraud</code>	Load credit card fraud dataset.
<code>load_wine</code>	Load wine dataset.
<code>load_breast_cancer</code>	Load breast cancer dataset.
<code>load_diabetes</code>	Load diabetes dataset.

evalml.demos.load_fraud

`evalml.demos.load_fraud(n_rows=None)`
 Load credit card fraud dataset. Binary classification problem

Parameters `n_rows` (*int*) – number of rows to return

Returns `X, y`

Return type `pd.DataFrame, pd.Series`

evalml.demos.load_wine

`evalml.demos.load_wine()`
 Load wine dataset. Multiclass problem

Returns `X, y`

Return type `pd.DataFrame, pd.Series`

evalml.demos.load_breast_cancer

`evalml.demos.load_breast_cancer()`
 Load breast cancer dataset. Multiclass problem

Returns `X, y`

Return type `pd.DataFrame, pd.Series`

evalml.demos.load_diabetes

`evalml.demos.load_diabetes()`
 Load diabetes dataset. Regression problem

Returns `X, y`

Return type `pd.DataFrame, pd.Series`

Preprocessing

<code>load_data</code>	Load features and labels from file(s).
<code>split_data</code>	Splits data into train and test sets.

evalml.preprocessing.load_data

`evalml.preprocessing.load_data(path, index, label, n_rows=None, drop=None, verbose=True, **kwargs)`
 Load features and labels from file(s).

Parameters

- **path** (*str*) – path to file(s)
- **index** (*str*) – column for index

- **label** (*str*) – column for labels
- **n_rows** (*int*) – number of rows to return
- **drop** (*list*) – columns to drop
- **verbose** (*bool*) – whether to print information about features and labels

Returns features and labels

Return type `pd.DataFrame`, `pd.Series`

evalml.preprocessing.split_data

`evalml.preprocessing.split_data(X, y, regression=False, test_size=0.2, random_state=None)`
Splits data into train and test sets.

Parameters

- **X** (`pd.DataFrame` or `np.array`) – data of shape `[n_samples, n_features]`
- **y** (`pd.Series`) – labels of length `[n_samples]`
- **regression** (*bool*) – if true, do not use stratified split
- **test_size** (*float*) – percent of train set to holdout for testing
- **random_state** (*int*) – seed for the random number generator

Returns features and labels each split into train and test sets

Return type `pd.DataFrame`, `pd.DataFrame`, `pd.Series`, `pd.Series`

AutoML

<code>AutoClassificationSearch</code>	Automatic pipeline search class for classification problems
<code>AutoRegressionSearch</code>	Automatic pipeline search for regression problems

evalml.AutoClassificationSearch

```
class evalml.AutoClassificationSearch(objective=None, multiclass=False,
                                     max_pipelines=None, max_time=None, patience=None,
                                     tolerance=None, model_types=None, cv=None,
                                     tuner=None, detect_label_leakage=True,
                                     start_iteration_callback=None,
                                     add_result_callback=None, additional_objectives=None,
                                     random_state=0, verbose=True)
```

Automatic pipeline search class for classification problems

Methods

<code>__init__</code>	Automated classifier pipeline search
-----------------------	--------------------------------------

evalml.AutoClassificationSearch.__init__

```
AutoClassificationSearch.__init__(objective=None, multiclass=False,
                                  max_pipelines=None, max_time=None, pa-
                                  tience=None, tolerance=None, model_types=None,
                                  cv=None, tuner=None, detect_label_leakage=True,
                                  start_iteration_callback=None,
                                  add_result_callback=None, addi-
                                  tional_objectives=None, random_state=0, ver-
                                 bose=True)
```

Automated classifier pipeline search

Parameters

- **objective** (*Object*) – the objective to optimize
- **multiclass** (*bool*) – If True, expecting multiclass data. By default: False.
- **max_pipelines** (*int*) – Maximum number of pipelines to search. If max_pipelines and max_time is not set, then max_pipelines will default to max_pipelines of 5.
- **max_time** (*int, str*) – Maximum time to search for pipelines. This will not start a new pipeline search after the duration has elapsed. If it is an integer, then the time will be in seconds. For strings, time can be specified as seconds, minutes, or hours.
- **patience** (*int*) – Number of iterations without improvement to stop search early. Must be positive. If None, early stopping is disabled. Defaults to None.
- **tolerance** (*float*) – Minimum percentage difference to qualify as score improvement for early stopping. Only applicable if patience is not None. Defaults to None.
- **model_types** (*list*) – The model types to search. By default searches over all model_types. Run evalml.list_model_types(“classification”) to see options.
- **cv** – cross validation method to use. By default StratifiedKfold
- **tuner** – the tuner class to use. Defaults to scikit-optimize tuner
- **detect_label_leakage** (*bool*) – If True, check input features for label leakage and warn if found. Defaults to true.
- **start_iteration_callback** (*callable*) – function called before each pipeline training iteration. Passed two parameters: pipeline_class, parameters.
- **add_result_callback** (*callable*) – function called after each pipeline training iteration. Passed two parameters: results, trained_pipeline.
- **additional_objectives** (*list*) – Custom set of objectives to score on. Will override default objectives for problem type if not empty.
- **random_state** (*int*) – the random_state
- **verbose** (*boolean*) – If True, turn verbosity on. Defaults to True

Attributes

<code>best_pipeline</code>	Returns the best model found
<code>rankings</code>	Returns the rankings of the models searched

evalml.AutoRegressionSearch

```
class evalml.AutoRegressionSearch(objective=None, max_pipelines=None, max_time=None,
                                patience=None, tolerance=None, model_types=None,
                                cv=None, tuner=None, detect_label_leakage=True,
                                start_iteration_callback=None, add_result_callback=None,
                                additional_objectives=None, random_state=0, ver-
                               bose=True)
```

Automatic pipeline search for regression problems

Methods

<code>__init__</code>	Automated regressors pipeline search
-----------------------	--------------------------------------

evalml.AutoRegressionSearch.__init__

```
AutoRegressionSearch.__init__(objective=None, max_pipelines=None, max_time=None,
                              patience=None, tolerance=None, model_types=None,
                              cv=None, tuner=None, detect_label_leakage=True,
                              start_iteration_callback=None, add_result_callback=None,
                              additional_objectives=None, random_state=0, ver-
                             bose=True)
```

Automated regressors pipeline search

Parameters

- **objective** (*Object*) – the objective to optimize
- **max_pipelines** (*int*) – Maximum number of pipelines to search. If max_pipelines and max_time is not set, then max_pipelines will default to max_pipelines of 5.
- **max_time** (*int, str*) – Maximum time to search for pipelines. This will not start a new pipeline search after the duration has elapsed. If it is an integer, then the time will be in seconds. For strings, time can be specified as seconds, minutes, or hours.
- **model_types** (*list*) – The model types to search. By default searches over all model_types. Run `evalml.list_model_types("regression")` to see options.
- **patience** (*int*) – Number of iterations without improvement to stop search early. Must be positive. If None, early stopping is disabled. Defaults to None.
- **tolerance** (*float*) – Minimum percentage difference to qualify as score improvement for early stopping. Only applicable if patience is not None. Defaults to None.
- **cv** – cross validation method to use. By default StratifiedKfold
- **tuner** – the tuner class to use. Defaults to scikit-optimize tuner
- **detect_label_leakage** (*bool*) – If True, check input features for label leakage and warn if found. Defaults to true.
- **start_iteration_callback** (*callable*) – function called before each pipeline training iteration. Passed two parameters: pipeline_class, parameters.
- **add_result_callback** (*callable*) – function called after each pipeline training iteration. Passed two parameters: results, trained_pipeline.

- **additional_objectives** (*list*) – Custom set of objectives to score on. Will override default objectives for problem type if not empty.
- **random_state** (*int*) – the random_state
- **verbose** (*boolean*) – If True, turn verbosity on. Defaults to True

Attributes

<code>best_pipeline</code>	Returns the best model found
<code>rankings</code>	Returns the rankings of the models searched

Plotting

<code>AutoClassificationSearch.plot.get_roc_data</code>	Gets data that can be used to create a ROC plot.
<code>AutoClassificationSearch.plot.generate_roc_plot</code>	Generate Receiver Operating Characteristic (ROC) plot for a given pipeline using cross-validation using the data returned from <code>get_roc_data()</code> .
<code>AutoRegressionSearch.plot.get_roc_data</code>	Gets data that can be used to create a ROC plot.
<code>AutoRegressionSearch.plot.generate_roc_plot</code>	Generate Receiver Operating Characteristic (ROC) plot for a given pipeline using cross-validation using the data returned from <code>get_roc_data()</code> .
<code>AutoClassificationSearch.plot.get_confusion_matrix_data</code>	Gets data that can be used to create a confusion matrix plot.
<code>AutoClassificationSearch.plot.generate_confusion_matrix</code>	Generate confusion matrix plot for a given pipeline using the data returned from <code>get_confusion_matrix_data()</code> .
<code>AutoRegressionSearch.plot.get_confusion_matrix_data</code>	Gets data that can be used to create a confusion matrix plot.
<code>AutoRegressionSearch.plot.generate_confusion_matrix</code>	Generate confusion matrix plot for a given pipeline using the data returned from <code>get_confusion_matrix_data()</code> .

evalml.AutoClassificationSearch.plot.get_roc_data

`AutoClassificationSearch.plot.get_roc_data` (*pipeline_id*)

Gets data that can be used to create a ROC plot.

Returns Dictionary containing metrics used to generate an ROC plot.

evalml.AutoClassificationSearch.plot.generate_roc_plot

`AutoClassificationSearch.plot.generate_roc_plot` (*pipeline_id*)

Generate Receiver Operating Characteristic (ROC) plot for a given pipeline using cross-validation using the data returned from `get_roc_data()`.

Returns `plotly`.Figure representing the ROC plot generated

evalml.AutoRegressionSearch.plot.get_roc_data

`AutoRegressionSearch.plot.get_roc_data(pipeline_id)`

Gets data that can be used to create a ROC plot.

Returns Dictionary containing metrics used to generate an ROC plot.

evalml.AutoRegressionSearch.plot.generate_roc_plot

`AutoRegressionSearch.plot.generate_roc_plot(pipeline_id)`

Generate Receiver Operating Characteristic (ROC) plot for a given pipeline using cross-validation using the data returned from `get_roc_data()`.

Returns `plotly.Figure` representing the ROC plot generated

evalml.AutoClassificationSearch.plot.get_confusion_matrix_data

`AutoClassificationSearch.plot.get_confusion_matrix_data(pipeline_id)`

Gets data that can be used to create a confusion matrix plot.

Returns List containing information used to generate a confusion matrix plot. Each element in the list contains the confusion matrix data for that fold.

evalml.AutoClassificationSearch.plot.generate_confusion_matrix

`AutoClassificationSearch.plot.generate_confusion_matrix(pipeline_id,
fold_num=None)`

Generate confusion matrix plot for a given pipeline using the data returned from `get_confusion_matrix_data()`.

Returns `plotly.Figure` representing the confusion matrix plot generated

evalml.AutoRegressionSearch.plot.get_confusion_matrix_data

`AutoRegressionSearch.plot.get_confusion_matrix_data(pipeline_id)`

Gets data that can be used to create a confusion matrix plot.

Returns List containing information used to generate a confusion matrix plot. Each element in the list contains the confusion matrix data for that fold.

evalml.AutoRegressionSearch.plot.generate_confusion_matrix

`AutoRegressionSearch.plot.generate_confusion_matrix(pipeline_id, fold_num=None)`

Generate confusion matrix plot for a given pipeline using the data returned from `get_confusion_matrix_data()`.

Returns `plotly.Figure` representing the confusion matrix plot generated

Model Types

list_model_types

List model type for a particular problem type

evalml.list_model_types

`evalml.list_model_types(problem_type)`

List model type for a particular problem type

Parameters `problem_types` (`ProblemTypes` or `str`) – binary, multiclass, or regression

Returns `model_types`, list of model types

Components

Transformers

<i>OneHotEncoder</i>	Creates one-hot encoding for non-numeric data
<i>RFRegressorSelectFromModel</i>	Selects top features based on importance weights using a Random Forest regressor
<i>RFClassifierSelectFromModel</i>	Selects top features based on importance weights using a Random Forest classifier
<i>SimpleImputer</i>	Imputes missing data with either mean, median and most_frequent for numerical data or most_frequent for categorical data
<i>StandardScaler</i>	Standardize features: removes mean and scales to unit variance

evalml.pipelines.components.OneHotEncoder

class `evalml.pipelines.components.OneHotEncoder`

Creates one-hot encoding for non-numeric data

Methods

<i>__init__</i>	Initialize self.
<i>describe</i>	Describe a component and its parameters
<i>fit</i>	Fits component to data
<i>fit_transform</i>	Fits on X and transforms X
<i>get_feature_names</i>	Returns names of transformed and added columns
<i>transform</i>	Transforms data X

evalml.pipelines.components.OneHotEncoder.__init__

`OneHotEncoder.__init__()`

Initialize self. See `help(type(self))` for accurate signature.

evalml.pipelines.components.OneHotEncoder.describe

`OneHotEncoder.describe(print_name=False, return_dict=False)`

Describe a component and its parameters

Parameters

- **print_name** (*bool, optional*) – whether to print name of component
- **return_dict** (*bool, optional*) – whether to return description as dictionary in the format {"name": name, "parameters": parameters}

Returns prints and returns dictionary

Return type None or dict

evalml.pipelines.components.OneHotEncoder.fit

`OneHotEncoder.fit(X, y=None)`

Fits component to data

Parameters

- **X** (*pd.DataFrame or np.array*) – the input training data of shape [n_samples, n_features]
- **y** (*pd.Series, optional*) – the target training labels of length [n_samples]

Returns self

evalml.pipelines.components.OneHotEncoder.fit_transform

`OneHotEncoder.fit_transform(X, y=None)`

Fits on X and transforms X

Parameters **X** (*pd.DataFrame*) – Data to fit and transform

Returns Transformed X

Return type pd.DataFrame

evalml.pipelines.components.OneHotEncoder.get_feature_names

`OneHotEncoder.get_feature_names()`

Returns names of transformed and added columns

Returns list of feature names not including dropped features

Return type list

evalml.pipelines.components.OneHotEncoder.transform

`OneHotEncoder.transform(X)`

Transforms data X

Parameters **X** (*pd.DataFrame*) – Data to transform

Returns Transformed X

Return type pd.DataFrame

evalml.pipelines.components.RFRegressorSelectFromModel

```
class evalml.pipelines.components.RFRegressorSelectFromModel (number_features=None,
                                                                n_estimators=10,
                                                                max_depth=None,
                                                                per-
                                                                cent_features=0.5,
                                                                threshold=-inf,
                                                                n_jobs=-1,    ran-
                                                                dom_state=0)
```

Selects top features based on importance weights using a Random Forest regressor

Methods

<code>__init__</code>	Initialize self.
<code>describe</code>	Describe a component and its parameters
<code>fit</code>	Fits component to data
<code>fit_transform</code>	Fits on X and transforms X
<code>get_indices</code>	Get integer index of features selected
<code>get_names</code>	Get names of selected features.
<code>transform</code>	Transforms data X

evalml.pipelines.components.RFRegressorSelectFromModel.__init__

```
RFRegressorSelectFromModel.__init__ (number_features=None,      n_estimators=10,
                                       max_depth=None,           percent_features=0.5,
                                       threshold=-inf, n_jobs=-1, random_state=0)
```

Initialize self. See help(type(self)) for accurate signature.

evalml.pipelines.components.RFRegressorSelectFromModel.describe

```
RFRegressorSelectFromModel.describe (print_name=False, return_dict=False)
```

Describe a component and its parameters

Parameters

- **print_name** (*bool, optional*) – whether to print name of component
- **return_dict** (*bool, optional*) – whether to return description as dictionary in the format {"name": name, "parameters": parameters}

Returns prints and returns dictionary

Return type None or dict

evalml.pipelines.components.RFRegressorSelectFromModel.fit

```
RFRegressorSelectFromModel.fit (X, y=None)
```

Fits component to data

Parameters

- **X** (*pd.DataFrame* or *np.array*) – the input training data of shape [n_samples, n_features]
- **y** (*pd.Series*, *optional*) – the target training labels of length [n_samples]

Returns self

`evalml.pipelines.components.RFRegressorSelectFromModel.fit_transform`

`RFRegressorSelectFromModel.fit_transform(X, y=None)`

Fits on X and transforms X

Parameters **X** (*pd.DataFrame*) – Data to fit and transform

Returns Transformed X

Return type *pd.DataFrame*

`evalml.pipelines.components.RFRegressorSelectFromModel.get_indices`

`RFRegressorSelectFromModel.get_indices()`

Get integer index of features selected

Returns list of indices

Return type list

`evalml.pipelines.components.RFRegressorSelectFromModel.get_names`

`RFRegressorSelectFromModel.get_names()`

Get names of selected features.

Returns list of the names of features selected

`evalml.pipelines.components.RFRegressorSelectFromModel.transform`

`RFRegressorSelectFromModel.transform(X)`

Transforms data X

Parameters **X** (*pd.DataFrame*) – Data to transform

Returns Transformed X

Return type *pd.DataFrame*

evalml.pipelines.components.RFClassifierSelectFromModel

```
class evalml.pipelines.components.RFClassifierSelectFromModel (number_features=None,  
                                                             n_estimators=10,  
                                                             max_depth=None,  
                                                             per-  
                                                             cent_features=0.5,  
                                                             threshold=-inf,  
                                                             n_jobs=-1, ran-  
                                                             dom_state=0)
```

Selects top features based on importance weights using a Random Forest classifier

Methods

<code>__init__</code>	Initialize self.
<code>describe</code>	Describe a component and its parameters
<code>fit</code>	Fits component to data
<code>fit_transform</code>	Fits on X and transforms X
<code>get_indices</code>	Get integer index of features selected
<code>get_names</code>	Get names of selected features.
<code>transform</code>	Transforms data X

evalml.pipelines.components.RFClassifierSelectFromModel.__init__

```
RFClassifierSelectFromModel.__init__ (number_features=None,      n_estimators=10,  
                                       max_depth=None,         percent_features=0.5,  
                                       threshold=-inf, n_jobs=-1, random_state=0)
```

Initialize self. See help(type(self)) for accurate signature.

evalml.pipelines.components.RFClassifierSelectFromModel.describe

```
RFClassifierSelectFromModel.describe (print_name=False, return_dict=False)
```

Describe a component and its parameters

Parameters

- **print_name** (*bool, optional*) – whether to print name of component
- **return_dict** (*bool, optional*) – whether to return description as dictionary in the format {"name": name, "parameters": parameters}

Returns prints and returns dictionary

Return type None or dict

evalml.pipelines.components.RFClassifierSelectFromModel.fit

```
RFClassifierSelectFromModel.fit (X, y=None)
```

Fits component to data

Parameters

- **X** (*pd.DataFrame* or *np.array*) – the input training data of shape [n_samples, n_features]
- **y** (*pd.Series*, *optional*) – the target training labels of length [n_samples]

Returns self

evalml.pipelines.components.RFClassifierSelectFromModel.fit_transform

`RFClassifierSelectFromModel.fit_transform(X, y=None)`

Fits on X and transforms X

Parameters **X** (*pd.DataFrame*) – Data to fit and transform

Returns Transformed X

Return type *pd.DataFrame*

evalml.pipelines.components.RFClassifierSelectFromModel.get_indices

`RFClassifierSelectFromModel.get_indices()`

Get integer index of features selected

Returns list of indices

Return type list

evalml.pipelines.components.RFClassifierSelectFromModel.get_names

`RFClassifierSelectFromModel.get_names()`

Get names of selected features.

Returns list of the names of features selected

evalml.pipelines.components.RFClassifierSelectFromModel.transform

`RFClassifierSelectFromModel.transform(X)`

Transforms data X

Parameters **X** (*pd.DataFrame*) – Data to transform

Returns Transformed X

Return type *pd.DataFrame*

evalml.pipelines.components.SimpleImputer

class `evalml.pipelines.components.SimpleImputer(impute_strategy='most_frequent')`

Imputes missing data with either mean, median and most_frequent for numerical data or most_frequent for categorical data

Methods

<code>__init__</code>	Initialize self.
<code>describe</code>	Describe a component and its parameters
<code>fit</code>	Fits component to data
<code>fit_transform</code>	Fits on X and transforms X
<code>transform</code>	Transforms data X

`evalml.pipelines.components.SimpleImputer.__init__`

`SimpleImputer.__init__` (*impute_strategy='most_frequent'*)
Initialize self. See `help(type(self))` for accurate signature.

`evalml.pipelines.components.SimpleImputer.describe`

`SimpleImputer.describe` (*print_name=False, return_dict=False*)
Describe a component and its parameters

Parameters

- **print_name** (*bool, optional*) – whether to print name of component
- **return_dict** (*bool, optional*) – whether to return description as dictionary in the format {"name": name, "parameters": parameters}

Returns prints and returns dictionary

Return type None or dict

`evalml.pipelines.components.SimpleImputer.fit`

`SimpleImputer.fit` (*X, y=None*)
Fits component to data

Parameters

- **X** (*pd.DataFrame or np.array*) – the input training data of shape [n_samples, n_features]
- **y** (*pd.Series, optional*) – the target training labels of length [n_samples]

Returns self

`evalml.pipelines.components.SimpleImputer.fit_transform`

`SimpleImputer.fit_transform` (*X, y=None*)
Fits on X and transforms X

Parameters **X** (*pd.DataFrame*) – Data to fit and transform

Returns Transformed X

Return type `pd.DataFrame`

evalml.pipelines.components.SimpleImputer.transform`SimpleImputer.transform(X)`

Transforms data X

Parameters `X` (`pd.DataFrame`) – Data to transform**Returns** Transformed X**Return type** `pd.DataFrame`**evalml.pipelines.components.StandardScaler****class** `evalml.pipelines.components.StandardScaler`

Standardize features: removes mean and scales to unit variance

Methods

<code>__init__</code>	Initialize self.
<code>describe</code>	Describe a component and its parameters
<code>fit</code>	Fits component to data
<code>fit_transform</code>	Fits on X and transforms X
<code>transform</code>	Transforms data X

evalml.pipelines.components.StandardScaler.__init__`StandardScaler.__init__()`Initialize self. See `help(type(self))` for accurate signature.**evalml.pipelines.components.StandardScaler.describe**`StandardScaler.describe(print_name=False, return_dict=False)`

Describe a component and its parameters

Parameters

- **print_name** (*bool, optional*) – whether to print name of component
- **return_dict** (*bool, optional*) – whether to return description as dictionary in the format {"name": name, "parameters": parameters}

Returns prints and returns dictionary**Return type** None or dict**evalml.pipelines.components.StandardScaler.fit**`StandardScaler.fit(X, y=None)`

Fits component to data

Parameters

- **X** (*pd.DataFrame* or *np.array*) – the input training data of shape [n_samples, n_features]
- **y** (*pd.Series*, optional) – the target training labels of length [n_samples]

Returns self

evalml.pipelines.components.StandardScaler.fit_transform

`StandardScaler.fit_transform(X, y=None)`

Fits on X and transforms X

Parameters **X** (*pd.DataFrame*) – Data to fit and transform

Returns Transformed X

Return type *pd.DataFrame*

evalml.pipelines.components.StandardScaler.transform

`StandardScaler.transform(X)`

Transforms data X

Parameters **X** (*pd.DataFrame*) – Data to transform

Returns Transformed X

Return type *pd.DataFrame*

Estimators

<i>LogisticRegressionClassifier</i>	Logistic Regression Classifier
<i>RandomForestClassifier</i>	Random Forest Classifier
<i>XGBoostClassifier</i>	XGBoost Classifier
<i>LinearRegressor</i>	Linear Regressor
<i>RandomForestRegressor</i>	Random Forest Regressor

evalml.pipelines.components.LogisticRegressionClassifier

class evalml.pipelines.components.**LogisticRegressionClassifier** (*penalty='l2', C=1.0, n_jobs=-1, random_state=0*)

Logistic Regression Classifier

Methods

<i>__init__</i>	Initialize self.
<i>describe</i>	Describe a component and its parameters
<i>fit</i>	Fits component to data
<i>predict</i>	Make predictions using selected features.

Continued on next page

Table 17 – continued from previous page

predict_proba

Make probability estimates for labels.

evalml.pipelines.components.LogisticRegressionClassifier.__init__

`LogisticRegressionClassifier.__init__(penalty='l2', C=1.0, n_jobs=-1, random_state=0)`
 Initialize self. See help(type(self)) for accurate signature.

evalml.pipelines.components.LogisticRegressionClassifier.describe

`LogisticRegressionClassifier.describe(print_name=False, return_dict=False)`
 Describe a component and its parameters

Parameters

- **print_name** (*bool, optional*) – whether to print name of component
- **return_dict** (*bool, optional*) – whether to return description as dictionary in the format {"name": name, "parameters": parameters}

Returns prints and returns dictionary**Return type** None or dict**evalml.pipelines.components.LogisticRegressionClassifier.fit**

`LogisticRegressionClassifier.fit(X, y=None)`
 Fits component to data

Parameters

- **X** (*pd.DataFrame or np.array*) – the input training data of shape [n_samples, n_features]
- **y** (*pd.Series, optional*) – the target training labels of length [n_samples]

Returns self**evalml.pipelines.components.LogisticRegressionClassifier.predict**

`LogisticRegressionClassifier.predict(X)`
 Make predictions using selected features.

Parameters **X** (*pd.DataFrame*) – features**Returns** estimated labels**Return type** pd.Series**evalml.pipelines.components.LogisticRegressionClassifier.predict_proba**

`LogisticRegressionClassifier.predict_proba(X)`
 Make probability estimates for labels.

Parameters **X** (*pd.DataFrame*) – features

Returns probability estimates

Return type pd.DataFrame

evalml.pipelines.components.RandomForestClassifier

```
class evalml.pipelines.components.RandomForestClassifier (n_estimators=10,  
max_depth=None,  
n_jobs=-1, ran-  
dom_state=0)
```

Random Forest Classifier

Methods

<code>__init__</code>	Initialize self.
<code>describe</code>	Describe a component and its parameters
<code>fit</code>	Fits component to data
<code>predict</code>	Make predictions using selected features.
<code>predict_proba</code>	Make probability estimates for labels.

evalml.pipelines.components.RandomForestClassifier.__init__

`RandomForestClassifier.__init__` (*n_estimators=10, max_depth=None, n_jobs=-1, ran-*
dom_state=0)
Initialize self. See help(type(self)) for accurate signature.

evalml.pipelines.components.RandomForestClassifier.describe

`RandomForestClassifier.describe` (*print_name=False, return_dict=False*)
Describe a component and its parameters

Parameters

- **print_name** (*bool, optional*) – whether to print name of component
- **return_dict** (*bool, optional*) – whether to return description as dictionary in the format {"name": name, "parameters": parameters}

Returns prints and returns dictionary

Return type None or dict

evalml.pipelines.components.RandomForestClassifier.fit

`RandomForestClassifier.fit` (*X, y=None*)
Fits component to data

Parameters

- **X** (*pd.DataFrame or np.array*) – the input training data of shape [n_samples, n_features]
- **y** (*pd.Series, optional*) – the target training labels of length [n_samples]

Returns self

evalml.pipelines.components.RandomForestClassifier.predict

RandomForestClassifier.**predict**(X)

Make predictions using selected features.

Parameters X (*pd.DataFrame*) – features

Returns estimated labels

Return type pd.Series

evalml.pipelines.components.RandomForestClassifier.predict_proba

RandomForestClassifier.**predict_proba**(X)

Make probability estimates for labels.

Parameters X (*pd.DataFrame*) – features

Returns probability estimates

Return type pd.DataFrame

evalml.pipelines.components.XGBoostClassifier

```
class evalml.pipelines.components.XGBoostClassifier(eta=0.1, max_depth=3,
                                                    min_child_weight=1, random_state=0)
```

XGBoost Classifier

Methods

<code>__init__</code>	Initialize self.
<code>describe</code>	Describe a component and its parameters
<code>fit</code>	Fits component to data
<code>predict</code>	Make predictions using selected features.
<code>predict_proba</code>	Make probability estimates for labels.

evalml.pipelines.components.XGBoostClassifier.__init__

XGBoostClassifier.**__init__**(eta=0.1, max_depth=3, min_child_weight=1, random_state=0)

Initialize self. See help(type(self)) for accurate signature.

evalml.pipelines.components.XGBoostClassifier.describe

XGBoostClassifier.**describe**(print_name=False, return_dict=False)

Describe a component and its parameters

Parameters

- **print_name** (*bool*, *optional*) – whether to print name of component

- **return_dict** (*bool, optional*) – whether to return description as dictionary in the format {"name": name, "parameters": parameters}

Returns prints and returns dictionary

Return type None or dict

evalml.pipelines.components.XGBoostClassifier.fit

XGBoostClassifier.**fit** (*X, y=None*)

Fits component to data

Parameters

- **X** (*pd.DataFrame or np.array*) – the input training data of shape [n_samples, n_features]
- **y** (*pd.Series, optional*) – the target training labels of length [n_samples]

Returns self

evalml.pipelines.components.XGBoostClassifier.predict

XGBoostClassifier.**predict** (*X*)

Make predictions using selected features.

Parameters **X** (*pd.DataFrame*) – features

Returns estimated labels

Return type pd.Series

evalml.pipelines.components.XGBoostClassifier.predict_proba

XGBoostClassifier.**predict_proba** (*X*)

Make probability estimates for labels.

Parameters **X** (*pd.DataFrame*) – features

Returns probability estimates

Return type pd.DataFrame

evalml.pipelines.components.LinearRegressor

class evalml.pipelines.components.**LinearRegressor** (*fit_intercept=True, normalize=False, n_jobs=-1*)

Linear Regressor

Methods

<code>__init__</code>	Initialize self.
<code>describe</code>	Describe a component and its parameters
<code>fit</code>	Fits component to data

Continued on next page

Table 20 – continued from previous page

<code>predict</code>	Make predictions using selected features.
<code>predict_proba</code>	Make probability estimates for labels.

evalml.pipelines.components.LinearRegressor.__init__

`LinearRegressor.__init__` (*fit_intercept=True, normalize=False, n_jobs=-1*)
 Initialize self. See help(type(self)) for accurate signature.

evalml.pipelines.components.LinearRegressor.describe

`LinearRegressor.describe` (*print_name=False, return_dict=False*)
 Describe a component and its parameters

Parameters

- **print_name** (*bool, optional*) – whether to print name of component
- **return_dict** (*bool, optional*) – whether to return description as dictionary in the format {"name": name, "parameters": parameters}

Returns prints and returns dictionary

Return type None or dict

evalml.pipelines.components.LinearRegressor.fit

`LinearRegressor.fit` (*X, y=None*)
 Fits component to data

Parameters

- **X** (*pd.DataFrame or np.array*) – the input training data of shape [n_samples, n_features]
- **y** (*pd.Series, optional*) – the target training labels of length [n_samples]

Returns self

evalml.pipelines.components.LinearRegressor.predict

`LinearRegressor.predict` (*X*)
 Make predictions using selected features.

Parameters **X** (*pd.DataFrame*) – features

Returns estimated labels

Return type pd.Series

evalml.pipelines.components.LinearRegressor.predict_proba

`LinearRegressor.predict_proba` (*X*)
 Make probability estimates for labels.

Parameters **X** (*pd.DataFrame*) – features

Returns probability estimates

Return type pd.DataFrame

evalml.pipelines.components.RandomForestRegressor

```
class evalml.pipelines.components.RandomForestRegressor (n_estimators=10,  
                                                         max_depth=None,  
                                                         n_jobs=-1,           ran-  
                                                         dom_state=0)
```

Random Forest Regressor

Methods

<code>__init__</code>	Initialize self.
<code>describe</code>	Describe a component and its parameters
<code>fit</code>	Fits component to data
<code>predict</code>	Make predictions using selected features.
<code>predict_proba</code>	Make probability estimates for labels.

evalml.pipelines.components.RandomForestRegressor.__init__

RandomForestRegressor.**__init__** (*n_estimators=10, max_depth=None, n_jobs=-1, ran-*
dom_state=0)
 Initialize self. See help(type(self)) for accurate signature.

evalml.pipelines.components.RandomForestRegressor.describe

RandomForestRegressor.**describe** (*print_name=False, return_dict=False*)
 Describe a component and its parameters

Parameters

- **print_name** (*bool, optional*) – whether to print name of component
- **return_dict** (*bool, optional*) – whether to return description as dictionary in the format {"name": name, "parameters": parameters}

Returns prints and returns dictionary

Return type None or dict

evalml.pipelines.components.RandomForestRegressor.fit

RandomForestRegressor.**fit** (*X, y=None*)
 Fits component to data

Parameters

- **X** (*pd.DataFrame or np.array*) – the input training data of shape [n_samples, n_features]
- **y** (*pd.Series, optional*) – the target training labels of length [n_samples]

Returns self

evalml.pipelines.components.RandomForestRegressor.predict

RandomForestRegressor.**predict** (*X*)

Make predictions using selected features.

Parameters *X* (*pd.DataFrame*) – features

Returns estimated labels

Return type *pd.Series*

evalml.pipelines.components.RandomForestRegressor.predict_proba

RandomForestRegressor.**predict_proba** (*X*)

Make probability estimates for labels.

Parameters *X* (*pd.DataFrame*) – features

Returns probability estimates

Return type *pd.DataFrame*

Pipelines

<code>get_pipelines</code>	Returns potential pipelines by model type
<code>save_pipeline</code>	Saves pipeline at file path
<code>load_pipeline</code>	Loads pipeline at file path

evalml.pipelines.get_pipelines

evalml.pipelines.**get_pipelines** (*problem_type*, *model_types=None*)

Returns potential pipelines by model type

Parameters

- **problem_type** (*ProblemTypes* or *str*) – the problem type the pipelines work for.
- **model_types** (*list[ModelTypes* or *str]*) – model types to match. if none, return all pipelines

Returns pipelines, list of all pipeline

evalml.pipelines.save_pipeline

evalml.pipelines.**save_pipeline** (*pipeline*, *file_path*)

Saves pipeline at file path

Parameters *file_path* (*str*) – location to save file

Returns None

evalml.pipelines.load_pipeline

evalml.pipelines.load_pipeline(*file_path*)

Loads pipeline at file path

Parameters *file_path* (*str*) – location to load file

Returns Pipeline obj

<i>PipelineBase</i>	
<i>RFClassificationPipeline</i>	Random Forest Pipeline for both binary and multiclass classification
<i>XGBoostPipeline</i>	XGBoost Pipeline for both binary and multiclass classification
<i>LogisticRegressionPipeline</i>	Logistic Regression Pipeline for both binary and multiclass classification
<i>RFRegressionPipeline</i>	Random Forest Pipeline for regression problems
<i>LinearRegressionPipeline</i>	Linear Regression Pipeline for regression problems

evalml.pipelines.PipelineBase

class evalml.pipelines.PipelineBase(*objective*, *component_list*, *n_jobs*, *random_state*)

Methods

<i>__init__</i>	Machine learning pipeline made out of transformers and a estimator.
<i>describe</i>	Outputs pipeline details including component parameters
<i>fit</i>	Build a model
<i>get_component</i>	Returns component by name
<i>predict</i>	Make predictions using selected features.
<i>predict_proba</i>	Make probability estimates for labels.
<i>score</i>	Evaluate model performance on current and additional objectives

evalml.pipelines.PipelineBase.__init__

PipelineBase.__init__(*objective*, *component_list*, *n_jobs*, *random_state*)

Machine learning pipeline made out of transformers and a estimator.

Parameters

- **objective** (*Object*) – the objective to optimize
- **component_list** (*list*) – List of components in order
- **random_state** (*int*) – random seed/state
- **n_jobs** (*int*) – Number of jobs to run in parallel

evalml.pipelines.PipelineBase.describe

PipelineBase.**describe** (*return_dict=False*)

Outputs pipeline details including component parameters

Parameters **return_dict** (*bool*) – If True, return dictionary of information about pipeline.
Defaults to false

Returns dictionary of all component parameters if return_dict is True, else None

Return type dict

evalml.pipelines.PipelineBase.fit

PipelineBase.**fit** (*X, y, objective_fit_size=0.2*)

Build a model

Parameters

- **X** (*pd.DataFrame or np.array*) – the input training data of shape [n_samples, n_features]
- **y** (*pd.Series*) – the target training labels of length [n_samples]
- **feature_types** (*list, optional*) – list of feature types. either numeric or categorical. categorical features will automatically be encoded

Returns self

evalml.pipelines.PipelineBase.get_component

PipelineBase.**get_component** (*name*)

Returns component by name

Parameters **name** (*str*) – name of component

Returns component to return

Return type Component

evalml.pipelines.PipelineBase.predict

PipelineBase.**predict** (*X*)

Make predictions using selected features.

Parameters **X** (*pd.DataFrame or np.array*) – data of shape [n_samples, n_features]

Returns estimated labels

Return type pd.Series

evalml.pipelines.PipelineBase.predict_proba

PipelineBase.**predict_proba** (*X*)

Make probability estimates for labels.

Parameters **X** (*pd.DataFrame or np.array*) – data of shape [n_samples, n_features]

Returns probability estimates

Return type `pd.DataFrame`

`evalml.pipelines.PipelineBase.score`

`PipelineBase.score(X, y, other_objectives=None)`

Evaluate model performance on current and additional objectives

Parameters

- **X** (`pd.DataFrame` or `np.array`) – data of shape `[n_samples, n_features]`
- **y** (`pd.Series`) – true labels of length `[n_samples]`
- **other_objectives** (`list`) – list of other objectives to score

Returns score, ordered dictionary of other objective scores

Return type `float, dict`

`evalml.pipelines.RFClassificationPipeline`

```
class evalml.pipelines.RFClassificationPipeline(objective, n_estimators, max_depth,
                                                impute_strategy, percent_features,
                                                number_features, n_jobs=-1, ran-
                                                dom_state=0)
```

Random Forest Pipeline for both binary and multiclass classification

Methods

<code>__init__</code>	Machine learning pipeline made out of transformers and a estimator.
<code>describe</code>	Outputs pipeline details including component parameters
<code>fit</code>	Build a model
<code>get_component</code>	Returns component by name
<code>predict</code>	Make predictions using selected features.
<code>predict_proba</code>	Make probability estimates for labels.
<code>score</code>	Evaluate model performance on current and additional objectives

`evalml.pipelines.RFClassificationPipeline.__init__`

```
RFClassificationPipeline.__init__(objective, n_estimators, max_depth, impute_strategy,
                                  percent_features, number_features, n_jobs=-1, ran-
                                  dom_state=0)
```

Machine learning pipeline made out of transformers and a estimator.

Parameters

- **objective** (`Object`) – the objective to optimize
- **component_list** (`list`) – List of components in order
- **random_state** (`int`) – random seed/state

- **n_jobs** (*int*) – Number of jobs to run in parallel

evalml.pipelines.RFClassificationPipeline.describe

RFClassificationPipeline.**describe** (*return_dict=False*)

Outputs pipeline details including component parameters

Parameters **return_dict** (*bool*) – If True, return dictionary of information about pipeline.
Defaults to false

Returns dictionary of all component parameters if return_dict is True, else None

Return type dict

evalml.pipelines.RFClassificationPipeline.fit

RFClassificationPipeline.**fit** (*X, y, objective_fit_size=0.2*)

Build a model

Parameters

- **X** (*pd.DataFrame* or *np.array*) – the input training data of shape [n_samples, n_features]
- **y** (*pd.Series*) – the target training labels of length [n_samples]
- **feature_types** (*list, optional*) – list of feature types. either numeric or categorical. categorical features will automatically be encoded

Returns self

evalml.pipelines.RFClassificationPipeline.get_component

RFClassificationPipeline.**get_component** (*name*)

Returns component by name

Parameters **name** (*str*) – name of component

Returns component to return

Return type Component

evalml.pipelines.RFClassificationPipeline.predict

RFClassificationPipeline.**predict** (*X*)

Make predictions using selected features.

Parameters **X** (*pd.DataFrame* or *np.array*) – data of shape [n_samples, n_features]

Returns estimated labels

Return type pd.Series

evalml.pipelines.RFClassificationPipeline.predict_proba**RFClassificationPipeline.predict_proba** (*X*)

Make probability estimates for labels.

Parameters *X* (*pd.DataFrame* or *np.array*) – data of shape [n_samples, n_features]**Returns** probability estimates**Return type** *pd.DataFrame***evalml.pipelines.RFClassificationPipeline.score****RFClassificationPipeline.score** (*X*, *y*, *other_objectives=None*)

Evaluate model performance on current and additional objectives

Parameters

- *X* (*pd.DataFrame* or *np.array*) – data of shape [n_samples, n_features]
- *y* (*pd.Series*) – true labels of length [n_samples]
- **other_objectives** (*list*) – list of other objectives to score

Returns score, ordered dictionary of other objective scores**Return type** float, dict**evalml.pipelines.XGBoostPipeline****class** **evalml.pipelines.XGBoostPipeline** (*objective*, *eta*, *min_child_weight*, *max_depth*, *impute_strategy*, *percent_features*, *number_features*, *n_estimators=10*, *n_jobs=-1*, *random_state=0*)

XGBoost Pipeline for both binary and multiclass classification

Methods

<code>__init__</code>	Machine learning pipeline made out of transformers and a estimator.
<code>describe</code>	Outputs pipeline details including component parameters
<code>fit</code>	Build a model
<code>get_component</code>	Returns component by name
<code>predict</code>	Make predictions using selected features.
<code>predict_proba</code>	Make probability estimates for labels.
<code>score</code>	Evaluate model performance on current and additional objectives

evalml.pipelines.XGBoostPipeline.__init__**XGBoostPipeline.__init__** (*objective*, *eta*, *min_child_weight*, *max_depth*, *impute_strategy*, *percent_features*, *number_features*, *n_estimators=10*, *n_jobs=-1*, *random_state=0*)

Machine learning pipeline made out of transformers and a estimator.

Parameters

- **objective** (*Object*) – the objective to optimize
- **component_list** (*list*) – List of components in order
- **random_state** (*int*) – random seed/state
- **n_jobs** (*int*) – Number of jobs to run in parallel

evalml.pipelines.XGBoostPipeline.describe

XGBoostPipeline.**describe** (*return_dict=False*)

Outputs pipeline details including component parameters

Parameters **return_dict** (*bool*) – If True, return dictionary of information about pipeline. Defaults to false

Returns dictionary of all component parameters if return_dict is True, else None

Return type dict

evalml.pipelines.XGBoostPipeline.fit

XGBoostPipeline.**fit** (*X, y, objective_fit_size=0.2*)

Build a model

Parameters

- **X** (*pd.DataFrame or np.array*) – the input training data of shape [n_samples, n_features]
- **y** (*pd.Series*) – the target training labels of length [n_samples]
- **feature_types** (*list, optional*) – list of feature types. either numeric of categorical. categorical features will automatically be encoded

Returns self

evalml.pipelines.XGBoostPipeline.get_component

XGBoostPipeline.**get_component** (*name*)

Returns component by name

Parameters **name** (*str*) – name of component

Returns component to return

Return type Component

evalml.pipelines.XGBoostPipeline.predict

XGBoostPipeline.**predict** (*X*)

Make predictions using selected features.

Parameters **X** (*pd.DataFrame or np.array*) – data of shape [n_samples, n_features]

Returns estimated labels

Return type `pd.Series`

`evalml.pipelines.XGBoostPipeline.predict_proba`

`XGBoostPipeline.predict_proba(X)`

Make probability estimates for labels.

Parameters **X** (`pd.DataFrame` or `np.array`) – data of shape `[n_samples, n_features]`

Returns probability estimates

Return type `pd.DataFrame`

`evalml.pipelines.XGBoostPipeline.score`

`XGBoostPipeline.score(X, y, other_objectives=None)`

Evaluate model performance on current and additional objectives

Parameters

- **X** (`pd.DataFrame` or `np.array`) – data of shape `[n_samples, n_features]`
- **y** (`pd.Series`) – true labels of length `[n_samples]`
- **other_objectives** (`list`) – list of other objectives to score

Returns score, ordered dictionary of other objective scores

Return type `float`, `dict`

`evalml.pipelines.LogisticRegressionPipeline`

class `evalml.pipelines.LogisticRegressionPipeline(objective, penalty, C, impute_strategy, number_features, n_jobs=-1, random_state=0)`

Logistic Regression Pipeline for both binary and multiclass classification

Methods

<code>__init__</code>	Machine learning pipeline made out of transformers and a estimator.
<code>describe</code>	Outputs pipeline details including component parameters
<code>fit</code>	Build a model
<code>get_component</code>	Returns component by name
<code>predict</code>	Make predictions using selected features.
<code>predict_proba</code>	Make probability estimates for labels.
<code>score</code>	Evaluate model performance on current and additional objectives

evalml.pipelines.LogisticRegressionPipeline.__init__

LogisticRegressionPipeline.__init__(*objective*, *penalty*, *C*, *impute_strategy*, *number_features*, *n_jobs=-1*, *random_state=0*)

Machine learning pipeline made out of transformers and a estimator.

Parameters

- **objective** (*Object*) – the objective to optimize
- **component_list** (*list*) – List of components in order
- **random_state** (*int*) – random seed/state
- **n_jobs** (*int*) – Number of jobs to run in parallel

evalml.pipelines.LogisticRegressionPipeline.describe

LogisticRegressionPipeline.**describe**(*return_dict=False*)

Outputs pipeline details including component parameters

Parameters **return_dict** (*bool*) – If True, return dictionary of information about pipeline.
Defaults to false

Returns dictionary of all component parameters if return_dict is True, else None

Return type dict

evalml.pipelines.LogisticRegressionPipeline.fit

LogisticRegressionPipeline.**fit**(*X*, *y*, *objective_fit_size=0.2*)

Build a model

Parameters

- **X** (*pd.DataFrame* or *np.array*) – the input training data of shape [n_samples, n_features]
- **y** (*pd.Series*) – the target training labels of length [n_samples]
- **feature_types** (*list*, *optional*) – list of feature types. either numeric or categorical. categorical features will automatically be encoded

Returns self

evalml.pipelines.LogisticRegressionPipeline.get_component

LogisticRegressionPipeline.**get_component**(*name*)

Returns component by name

Parameters **name** (*str*) – name of component

Returns component to return

Return type Component

evalml.pipelines.LogisticRegressionPipeline.predict

LogisticRegressionPipeline.**predict** (*X*)

Make predictions using selected features.

Parameters *X* (*pd.DataFrame* or *np.array*) – data of shape [n_samples, n_features]

Returns estimated labels

Return type *pd.Series*

evalml.pipelines.LogisticRegressionPipeline.predict_proba

LogisticRegressionPipeline.**predict_proba** (*X*)

Make probability estimates for labels.

Parameters *X* (*pd.DataFrame* or *np.array*) – data of shape [n_samples, n_features]

Returns probability estimates

Return type *pd.DataFrame*

evalml.pipelines.LogisticRegressionPipeline.score

LogisticRegressionPipeline.**score** (*X*, *y*, *other_objectives=None*)

Evaluate model performance on current and additional objectives

Parameters

- *X* (*pd.DataFrame* or *np.array*) – data of shape [n_samples, n_features]
- *y* (*pd.Series*) – true labels of length [n_samples]
- **other_objectives** (*list*) – list of other objectives to score

Returns score, ordered dictionary of other objective scores

Return type float, dict

evalml.pipelines.RFRegressionPipeline

class evalml.pipelines.**RFRegressionPipeline** (*objective*, *n_estimators*, *max_depth*, *impute_strategy*, *percent_features*, *number_features*, *n_jobs=-1*, *random_state=0*)

Random Forest Pipeline for regression problems

Methods

<code>__init__</code>	Machine learning pipeline made out of transformers and a estimator.
<code>describe</code>	Outputs pipeline details including component parameters
<code>fit</code>	Build a model
<code>get_component</code>	Returns component by name

Continued on next page

Table 28 – continued from previous page

<code>predict</code>	Make predictions using selected features.
<code>predict_proba</code>	Make probability estimates for labels.
<code>score</code>	Evaluate model performance on current and additional objectives

evalml.pipelines.RFRegressionPipeline.__init__

`RFRegressionPipeline.__init__(objective, n_estimators, max_depth, impute_strategy, percent_features, number_features, n_jobs=-1, random_state=0)`

Machine learning pipeline made out of transformers and a estimator.

Parameters

- **objective** (*Object*) – the objective to optimize
- **component_list** (*list*) – List of components in order
- **random_state** (*int*) – random seed/state
- **n_jobs** (*int*) – Number of jobs to run in parallel

evalml.pipelines.RFRegressionPipeline.describe

`RFRegressionPipeline.describe(return_dict=False)`

Outputs pipeline details including component parameters

Parameters **return_dict** (*bool*) – If True, return dictionary of information about pipeline. Defaults to false

Returns dictionary of all component parameters if return_dict is True, else None

Return type dict

evalml.pipelines.RFRegressionPipeline.fit

`RFRegressionPipeline.fit(X, y, objective_fit_size=0.2)`

Build a model

Parameters

- **X** (*pd.DataFrame or np.array*) – the input training data of shape [n_samples, n_features]
- **y** (*pd.Series*) – the target training labels of length [n_samples]
- **feature_types** (*list, optional*) – list of feature types. either numeric or categorical. categorical features will automatically be encoded

Returns self

evalml.pipelines.RFRegressionPipeline.get_component

`RFRegressionPipeline.get_component(name)`

Returns component by name

Parameters **name** (*str*) – name of component

Returns component to return

Return type Component

`evalml.pipelines.RFRegressionPipeline.predict`

`RFRegressionPipeline.predict(X)`

Make predictions using selected features.

Parameters **X** (*pd.DataFrame* or *np.array*) – data of shape [n_samples, n_features]

Returns estimated labels

Return type *pd.Series*

`evalml.pipelines.RFRegressionPipeline.predict_proba`

`RFRegressionPipeline.predict_proba(X)`

Make probability estimates for labels.

Parameters **X** (*pd.DataFrame* or *np.array*) – data of shape [n_samples, n_features]

Returns probability estimates

Return type *pd.DataFrame*

`evalml.pipelines.RFRegressionPipeline.score`

`RFRegressionPipeline.score(X, y, other_objectives=None)`

Evaluate model performance on current and additional objectives

Parameters

- **X** (*pd.DataFrame* or *np.array*) – data of shape [n_samples, n_features]
- **y** (*pd.Series*) – true labels of length [n_samples]
- **other_objectives** (*list*) – list of other objectives to score

Returns score, ordered dictionary of other objective scores

Return type float, dict

`evalml.pipelines.LinearRegressionPipeline`

```
class evalml.pipelines.LinearRegressionPipeline(objective, random_state, num-  
ber_features, impute_strategy, nor-  
malize=False, fit_intercept=True,  
n_jobs=-1)
```

Linear Regression Pipeline for regression problems

Methods

<code>__init__</code>	Machine learning pipeline made out of transformers and a estimator.
<code>describe</code>	Outputs pipeline details including component parameters
<code>fit</code>	Build a model
<code>get_component</code>	Returns component by name
<code>predict</code>	Make predictions using selected features.
<code>predict_proba</code>	Make probability estimates for labels.
<code>score</code>	Evaluate model performance on current and additional objectives

`evalml.pipelines.LinearRegressionPipeline.__init__`

`LinearRegressionPipeline.__init__(objective, random_state, number_features, impute_strategy, normalize=False, fit_intercept=True, n_jobs=-1)`

Machine learning pipeline made out of transformers and a estimator.

Parameters

- **objective** (*Object*) – the objective to optimize
- **component_list** (*list*) – List of components in order
- **random_state** (*int*) – random seed/state
- **n_jobs** (*int*) – Number of jobs to run in parallel

`evalml.pipelines.LinearRegressionPipeline.describe`

`LinearRegressionPipeline.describe(return_dict=False)`

Outputs pipeline details including component parameters

Parameters **return_dict** (*bool*) – If True, return dictionary of information about pipeline. Defaults to false

Returns dictionary of all component parameters if return_dict is True, else None

Return type dict

`evalml.pipelines.LinearRegressionPipeline.fit`

`LinearRegressionPipeline.fit(X, y, objective_fit_size=0.2)`

Build a model

Parameters

- **X** (*pd.DataFrame or np.array*) – the input training data of shape [n_samples, n_features]
- **y** (*pd.Series*) – the target training labels of length [n_samples]
- **feature_types** (*list, optional*) – list of feature types. either numeric or categorical. categorical features will automatically be encoded

Returns self

evalml.pipelines.LinearRegressionPipeline.get_component

`LinearRegressionPipeline.get_component(name)`

Returns component by name

Parameters `name` (*str*) – name of component

Returns component to return

Return type Component

evalml.pipelines.LinearRegressionPipeline.predict

`LinearRegressionPipeline.predict(X)`

Make predictions using selected features.

Parameters `X` (*pd.DataFrame* or *np.array*) – data of shape [n_samples, n_features]

Returns estimated labels

Return type *pd.Series*

evalml.pipelines.LinearRegressionPipeline.predict_proba

`LinearRegressionPipeline.predict_proba(X)`

Make probability estimates for labels.

Parameters `X` (*pd.DataFrame* or *np.array*) – data of shape [n_samples, n_features]

Returns probability estimates

Return type *pd.DataFrame*

evalml.pipelines.LinearRegressionPipeline.score

`LinearRegressionPipeline.score(X, y, other_objectives=None)`

Evaluate model performance on current and additional objectives

Parameters

- `X` (*pd.DataFrame* or *np.array*) – data of shape [n_samples, n_features]
- `y` (*pd.Series*) – true labels of length [n_samples]
- `other_objectives` (*list*) – list of other objectives to score

Returns score, ordered dictionary of other objective scores

Return type float, dict

Plotting

PipelineBase.plot

alias of `evalml.pipelines.pipeline_plots.PipelinePlots`

evalml.pipelines.PipelineBase.plot

`pipelines.PipelineBase.plot` (*to_file=None*)

Create a graph of the pipeline, in a format similar to a UML diagram.

Parameters `to_file` (*str, optional*) – Path to where the plot should be saved. If set to None (as by default), the plot will not be saved.

Returns Graph object that can directly be displayed in Jupyter notebooks.

Return type `graphviz.Digraph`

Objective Functions

Domain Specific

<i>FraudCost</i>	Score the percentage of money lost of the total transaction amount process due to fraud
<i>LeadScoring</i>	Lead scoring

evalml.objectives.FraudCost

class `evalml.objectives.FraudCost` (*retry_percentage=0.5, interchange_fee=0.02, fraud_payout_percentage=1.0, amount_col='amount', verbose=False*)

Score the percentage of money lost of the total transaction amount process due to fraud

Methods

<i>__init__</i>	Create instance of FraudCost
<i>decision_function</i>	Determine if transaction is fraud given predicted probabilities, dataframe with transaction amount, and threshold
<i>fit</i>	Learn the objective function based on the predictions from a model.
<i>objective_function</i>	Calculate amount lost to fraud per transaction given predictions, true values, and dataframe with transaction amount
<i>predict</i>	Apply the learned objective function to the output of a model.
<i>score</i>	Calculate score from applying fitted objective to predicted values
<i>supports_problem_type</i>	Checks if objective supports given ProblemType

evalml.objectives.FraudCost.__init__

`FraudCost.__init__` (*retry_percentage=0.5, interchange_fee=0.02, fraud_payout_percentage=1.0, amount_col='amount', verbose=False*)

Create instance of FraudCost

Parameters

- **retry_percentage** (*float*) – what percentage of customers will retry a transaction if it is declined? Between 0 and 1. Defaults to .5
- **interchange_fee** (*float*) – how much of each successful transaction do you collect? Between 0 and 1. Defaults to .02
- **fraud_payout_percentage** (*float*) – how percentage of fraud will you be unable to collect. Between 0 and 1. Defaults to 1.0
- **amount_col** (*str*) – name of column in data that contains the amount. defaults to “amount”

evalml.objectives.FraudCost.decision_function

`FraudCost.decision_function(y_predicted, extra_cols, threshold)`

Determine if transaction is fraud given predicted probabilities, dataframe with transaction amount, and threshold

Parameters

- **y_predicted** (*pd.Series*) – predicted labels
- **extra_cols** (*pd.DataFrame*) – extra data needed
- **threshold** (*float*) – dollar threshold to determine if transaction is fraud

Returns series of predicted fraud label using extra cols and threshold

Return type *pd.Series*

evalml.objectives.FraudCost.fit

`FraudCost.fit(y_predicted, y_true, extra_cols=None)`

Learn the objective function based on the predictions from a model.

If `needs_fitting` is false, this method won't be called

Parameters

- **y_predicted** (*list*) – the predictions from the model. If `needs_proba` is True, it is the probability estimates
- **y_true** (*list*) – the ground truth for the predictions.
- **extra_cols** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if `uses_extra_columns` is True.

Returns *self*

evalml.objectives.FraudCost.objective_function

`FraudCost.objective_function(y_predicted, y_true, extra_cols)`

Calculate amount lost to fraud per transaction given predictions, true values, and dataframe with transaction amount

Parameters

- **y_predicted** (*pd.Series*) – predicted fraud labels
- **y_true** (*pd.Series*) – true fraud labels

- **extra_cols** (*pd.DataFrame*) – extra data needed

Returns amount lost to fraud per transaction

Return type float

evalml.objectives.FraudCost.predict

`FraudCost.predict(y_predicted, extra_cols=None)`

Apply the learned objective function to the output of a model.

If `needs_fitting` is false, this method won't be called

Parameters **y_predicted** – the prediction to transform to final prediction

Returns predictions

evalml.objectives.FraudCost.score

`FraudCost.score(y_predicted, y_true, extra_cols=None)`

Calculate score from applying fitted objective to predicted values

If a higher score is better than a lower score, set `greater_is_better` attribute to True

Parameters

- **y_predicted** (*list*) – the predictions from the model. If `needs_proba` is True, it is the probability estimates
- **y_true** (*list*) – the ground truth for the predictions.
- **extra_cols** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if `uses_extra_columns` is True.

Returns score

evalml.objectives.FraudCost.supports_problem_type

classmethod `FraudCost.supports_problem_type(problem_type)`

Checks if objective supports given `ProblemType`

Parameters **problem_type** (*str or ProblemType*) – problem type to check

Returns whether objective supports `ProblemType`

Return type bool

evalml.objectives.LeadScoring

class `evalml.objectives.LeadScoring(true_positives=1, false_positives=-1, verbose=False)`

Lead scoring

Methods

<code>__init__</code>	Create instance.
<code>decision_function</code>	
<code>fit</code>	Learn the objective function based on the predictions from a model.
<code>objective_function</code>	
<code>predict</code>	Apply the learned objective function to the output of a model.
<code>score</code>	Calculate score from applying fitted objective to predicted values
<code>supports_problem_type</code>	Checks if objective supports given ProblemType

`evalml.objectives.LeadScoring.__init__`

`LeadScoring.__init__(true_positives=1, false_positives=-1, verbose=False)`
Create instance.

Parameters

- **label** (*int*) – label to optimize threshold for
- **true_positives** (*int*) – reward for a true positive
- **false_positives** (*int*) – cost for a false positive. Should be negative.

`evalml.objectives.LeadScoring.decision_function`

`LeadScoring.decision_function(y_predicted, threshold)`

`evalml.objectives.LeadScoring.fit`

`LeadScoring.fit(y_predicted, y_true, extra_cols=None)`
Learn the objective function based on the predictions from a model.
If `needs_fitting` is false, this method won't be called

Parameters

- **y_predicted** (*list*) – the predictions from the model. If `needs_proba` is True, it is the probability estimates
- **y_true** (*list*) – the ground truth for the predictions.
- **extra_cols** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if `uses_extra_columns` is True.

Returns self

`evalml.objectives.LeadScoring.objective_function`

`LeadScoring.objective_function(y_predicted, y_true)`

evalml.objectives.LeadScoring.predict

`LeadScoring.predict` (*y_predicted*, *extra_cols=None*)

Apply the learned objective function to the output of a model.

If `needs_fitting` is false, this method won't be called

Parameters `y_predicted` – the prediction to transform to final prediction

Returns predictions

evalml.objectives.LeadScoring.score

`LeadScoring.score` (*y_predicted*, *y_true*, *extra_cols=None*)

Calculate score from applying fitted objective to predicted values

If a higher score is better than a lower score, set `greater_is_better` attribute to True

Parameters

- `y_predicted` (*list*) – the predictions from the model. If `needs_proba` is True, it is the probability estimates
- `y_true` (*list*) – the ground truth for the predictions.
- `extra_cols` (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if `uses_extra_columns` is True.

Returns score

evalml.objectives.LeadScoring.supports_problem_type

classmethod `LeadScoring.supports_problem_type` (*problem_type*)

Checks if objective supports given `ProblemType`

Parameters `problem_type` (*str* or *ProblemType*) – problem type to check

Returns whether objective supports `ProblemType`

Return type bool

Classification

<i>F1</i>	F1 score for binary classification
<i>F1Micro</i>	F1 score for multiclass classification using micro averaging
<i>F1Macro</i>	F1 score for multiclass classification using macro averaging
<i>F1Weighted</i>	F1 score for multiclass classification using weighted averaging
<i>Precision</i>	Precision score for binary classification
<i>PrecisionMicro</i>	Precision score for multiclass classification using micro averaging
<i>PrecisionMacro</i>	Precision score for multiclass classification using macro averaging

Continued on next page

Table 34 – continued from previous page

<i>PrecisionWeighted</i>	Precision score for multiclass classification using weighted averaging
<i>Recall</i>	Recall score for binary classification
<i>RecallMicro</i>	Recall score for multiclass classification using micro averaging
<i>RecallMacro</i>	Recall score for multiclass classification using macro averaging
<i>RecallWeighted</i>	Recall score for multiclass classification using weighted averaging
<i>AUC</i>	AUC score for binary classification
<i>AUCMicro</i>	AUC score for multiclass classification using micro averaging
<i>AUCMacro</i>	AUC score for multiclass classification using macro averaging
<i>AUCWeighted</i>	AUC Score for multiclass classification using weighted averaging
<i>LogLoss</i>	Log Loss for both binary and multiclass classification
<i>MCC</i>	Matthews correlation coefficient for both binary and multiclass classification
<i>ROC</i>	Receiver Operating Characteristic score for binary classification.
<i>ConfusionMatrix</i>	Confusion matrix for classification problems

evalml.objectives.F1

class evalml.objectives.F1 (*verbose=False*)
F1 score for binary classification

Methods

<i>__init__</i>	Initialize self.
<i>fit</i>	Learn the objective function based on the predictions from a model.
<i>predict</i>	Apply the learned objective function to the output of a model.
<i>score</i>	Calculate score from applying fitted objective to predicted values
<i>supports_problem_type</i>	Checks if objective supports given ProblemType

evalml.objectives.F1.__init__

F1.**__init__** (*verbose=False*)
Initialize self. See help(type(self)) for accurate signature.

evalml.objectives.F1.fit

F1.**fit** (*y_predicted, y_true, extra_cols=None*)
Learn the objective function based on the predictions from a model.

If `needs_fitting` is false, this method won't be called

Parameters

- **`y_predicted`** (*list*) – the predictions from the model. If `needs_proba` is True, it is the probability estimates
- **`y_true`** (*list*) – the ground truth for the predictions.
- **`extra_cols`** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if `uses_extra_columns` is True.

Returns `self`

`evalml.objectives.F1.predict`

`F1.predict` (*y_predicted*, *extra_cols=None*)

Apply the learned objective function to the output of a model.

If `needs_fitting` is false, this method won't be called

Parameters **`y_predicted`** – the prediction to transform to final prediction

Returns predictions

`evalml.objectives.F1.score`

`F1.score` (*y_predicted*, *y_true*)

Calculate score from applying fitted objective to predicted values

If a higher score is better than a lower score, set `greater_is_better` attribute to True

Parameters

- **`y_predicted`** (*list*) – the predictions from the model. If `needs_proba` is True, it is the probability estimates
- **`y_true`** (*list*) – the ground truth for the predictions.
- **`extra_cols`** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if `uses_extra_columns` is True.

Returns score

`evalml.objectives.F1.supports_problem_type`

`classmethod F1.supports_problem_type` (*problem_type*)

Checks if objective supports given `ProblemType`

Parameters **`problem_type`** (*str or ProblemType*) – problem type to check

Returns whether objective supports `ProblemType`

Return type bool

`evalml.objectives.F1Micro`

`class evalml.objectives.F1Micro` (*verbose=False*)

F1 score for multiclass classification using micro averaging

Methods

<code>__init__</code>	Initialize self.
<code>fit</code>	Learn the objective function based on the predictions from a model.
<code>predict</code>	Apply the learned objective function to the output of a model.
<code>score</code>	Calculate score from applying fitted objective to predicted values
<code>supports_problem_type</code>	Checks if objective supports given ProblemType

`evalml.objectives.F1Micro.__init__`

`F1Micro.__init__(verbose=False)`
Initialize self. See `help(type(self))` for accurate signature.

`evalml.objectives.F1Micro.fit`

`F1Micro.fit(y_predicted, y_true, extra_cols=None)`
Learn the objective function based on the predictions from a model.
If `needs_fitting` is false, this method won't be called

Parameters

- **`y_predicted`** (*list*) – the predictions from the model. If `needs_proba` is True, it is the probability estimates
- **`y_true`** (*list*) – the ground truth for the predictions.
- **`extra_cols`** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if `uses_extra_columns` is True.

Returns self

`evalml.objectives.F1Micro.predict`

`F1Micro.predict(y_predicted, extra_cols=None)`
Apply the learned objective function to the output of a model.
If `needs_fitting` is false, this method won't be called

Parameters **`y_predicted`** – the prediction to transform to final prediction

Returns predictions

`evalml.objectives.F1Micro.score`

`F1Micro.score(y_predicted, y_true)`
Calculate score from applying fitted objective to predicted values
If a higher score is better than a lower score, set `greater_is_better` attribute to True

Parameters

- **y_predicted** (*list*) – the predictions from the model. If `needs_proba` is `True`, it is the probability estimates
- **y_true** (*list*) – the ground truth for the predictions.
- **extra_cols** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if `uses_extra_columns` is `True`.

Returns score

`evalml.objectives.F1Micro.supports_problem_type`

classmethod `F1Micro.supports_problem_type` (*problem_type*)

Checks if objective supports given `ProblemType`

Parameters `problem_type` (*str or ProblemType*) – problem type to check

Returns whether objective supports `ProblemType`

Return type bool

`evalml.objectives.F1Macro`

class `evalml.objectives.F1Macro` (*verbose=False*)

F1 score for multiclass classification using macro averaging

Methods

<code>__init__</code>	Initialize self.
<code>fit</code>	Learn the objective function based on the predictions from a model.
<code>predict</code>	Apply the learned objective function to the output of a model.
<code>score</code>	Calculate score from applying fitted objective to predicted values
<code>supports_problem_type</code>	Checks if objective supports given <code>ProblemType</code>

`evalml.objectives.F1Macro.__init__`

`F1Macro.__init__` (*verbose=False*)

Initialize self. See `help(type(self))` for accurate signature.

`evalml.objectives.F1Macro.fit`

`F1Macro.fit` (*y_predicted, y_true, extra_cols=None*)

Learn the objective function based on the predictions from a model.

If `needs_fitting` is `false`, this method won't be called

Parameters

- **y_predicted** (*list*) – the predictions from the model. If `needs_proba` is `True`, it is the probability estimates

- **y_true** (*list*) – the ground truth for the predictions.
- **extra_cols** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if `uses_extra_columns` is `True`.

Returns `self`

evalml.objectives.F1Macro.predict

F1Macro.predict (*y_predicted*, *extra_cols=None*)

Apply the learned objective function to the output of a model.

If `needs_fitting` is `false`, this method won't be called

Parameters **y_predicted** – the prediction to transform to final prediction

Returns predictions

evalml.objectives.F1Macro.score

F1Macro.score (*y_predicted*, *y_true*)

Calculate score from applying fitted objective to predicted values

If a higher score is better than a lower score, set `greater_is_better` attribute to `True`

Parameters

- **y_predicted** (*list*) – the predictions from the model. If `needs_proba` is `True`, it is the probability estimates
- **y_true** (*list*) – the ground truth for the predictions.
- **extra_cols** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if `uses_extra_columns` is `True`.

Returns score

evalml.objectives.F1Macro.supports_problem_type

classmethod **F1Macro.supports_problem_type** (*problem_type*)

Checks if objective supports given `ProblemType`

Parameters **problem_type** (*str* or *ProblemType*) – problem type to check

Returns whether objective supports `ProblemType`

Return type `bool`

evalml.objectives.F1Weighted

class **evalml.objectives.F1Weighted** (*verbose=False*)

F1 score for multiclass classification using weighted averaging

Methods

<code>__init__</code>	Initialize self.
<code>fit</code>	Learn the objective function based on the predictions from a model.
<code>predict</code>	Apply the learned objective function to the output of a model.
<code>score</code>	Calculate score from applying fitted objective to predicted values
<code>supports_problem_type</code>	Checks if objective supports given ProblemType

`evalml.objectives.F1Weighted.__init__`

`F1Weighted.__init__(verbose=False)`
Initialize self. See `help(type(self))` for accurate signature.

`evalml.objectives.F1Weighted.fit`

`F1Weighted.fit(y_predicted, y_true, extra_cols=None)`
Learn the objective function based on the predictions from a model.
If `needs_fitting` is false, this method won't be called

Parameters

- **`y_predicted`** (*list*) – the predictions from the model. If `needs_proba` is True, it is the probability estimates
- **`y_true`** (*list*) – the ground truth for the predictions.
- **`extra_cols`** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if `uses_extra_columns` is True.

Returns self

`evalml.objectives.F1Weighted.predict`

`F1Weighted.predict(y_predicted, extra_cols=None)`
Apply the learned objective function to the output of a model.
If `needs_fitting` is false, this method won't be called

Parameters **`y_predicted`** – the prediction to transform to final prediction

Returns predictions

`evalml.objectives.F1Weighted.score`

`F1Weighted.score(y_predicted, y_true)`
Calculate score from applying fitted objective to predicted values
If a higher score is better than a lower score, set `greater_is_better` attribute to True

Parameters

- **`y_predicted`** (*list*) – the predictions from the model. If `needs_proba` is True, it is the probability estimates

- **y_true** (*list*) – the ground truth for the predictions.
- **extra_cols** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if `uses_extra_columns` is `True`.

Returns score

`evalml.objectives.F1Weighted.supports_problem_type`

classmethod `F1Weighted.supports_problem_type(problem_type)`

Checks if objective supports given `ProblemType`

Parameters **problem_type** (*str* or *ProblemType*) – problem type to check

Returns whether objective supports `ProblemType`

Return type `bool`

`evalml.objectives.Precision`

class `evalml.objectives.Precision(verbose=False)`

Precision score for binary classification

Methods

<code>__init__</code>	Initialize self.
<code>fit</code>	Learn the objective function based on the predictions from a model.
<code>predict</code>	Apply the learned objective function to the output of a model.
<code>score</code>	Calculate score from applying fitted objective to predicted values
<code>supports_problem_type</code>	Checks if objective supports given <code>ProblemType</code>

`evalml.objectives.Precision.__init__`

`Precision.__init__(verbose=False)`

Initialize self. See `help(type(self))` for accurate signature.

`evalml.objectives.Precision.fit`

`Precision.fit(y_predicted, y_true, extra_cols=None)`

Learn the objective function based on the predictions from a model.

If `needs_fitting` is `false`, this method won't be called

Parameters

- **y_predicted** (*list*) – the predictions from the model. If `needs_proba` is `True`, it is the probability estimates
- **y_true** (*list*) – the ground truth for the predictions.

- **extra_cols** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if `uses_extra_columns` is `True`.

Returns self

evalml.objectives.Precision.predict

`Precision.predict(y_predicted, extra_cols=None)`

Apply the learned objective function to the output of a model.

If `needs_fitting` is `false`, this method won't be called

Parameters `y_predicted` – the prediction to transform to final prediction

Returns predictions

evalml.objectives.Precision.score

`Precision.score(y_predicted, y_true)`

Calculate score from applying fitted objective to predicted values

If a higher score is better than a lower score, set `greater_is_better` attribute to `True`

Parameters

- **y_predicted** (*list*) – the predictions from the model. If `needs_proba` is `True`, it is the probability estimates
- **y_true** (*list*) – the ground truth for the predictions.
- **extra_cols** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if `uses_extra_columns` is `True`.

Returns score

evalml.objectives.Precision.supports_problem_type

classmethod `Precision.supports_problem_type(problem_type)`

Checks if objective supports given `ProblemType`

Parameters `problem_type` (*str or ProblemType*) – problem type to check

Returns whether objective supports `ProblemType`

Return type bool

evalml.objectives.PrecisionMicro

class `evalml.objectives.PrecisionMicro(verbose=False)`

Precision score for multiclass classification using micro averaging

Methods

`__init__`

Initialize self.

Continued on next page

Table 40 – continued from previous page

<i>fit</i>	Learn the objective function based on the predictions from a model.
<i>predict</i>	Apply the learned objective function to the output of a model.
<i>score</i>	Calculate score from applying fitted objective to predicted values
<i>supports_problem_type</i>	Checks if objective supports given ProblemType

evalml.objectives.PrecisionMicro.__init__

`PrecisionMicro.__init__(verbose=False)`
Initialize self. See `help(type(self))` for accurate signature.

evalml.objectives.PrecisionMicro.fit

`PrecisionMicro.fit(y_predicted, y_true, extra_cols=None)`
Learn the objective function based on the predictions from a model.
If `needs_fitting` is false, this method won't be called

Parameters

- **y_predicted** (*list*) – the predictions from the model. If `needs_proba` is True, it is the probability estimates
- **y_true** (*list*) – the ground truth for the predictions.
- **extra_cols** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if `uses_extra_columns` is True.

Returns self

evalml.objectives.PrecisionMicro.predict

`PrecisionMicro.predict(y_predicted, extra_cols=None)`
Apply the learned objective function to the output of a model.
If `needs_fitting` is false, this method won't be called

Parameters **y_predicted** – the prediction to transform to final prediction

Returns predictions

evalml.objectives.PrecisionMicro.score

`PrecisionMicro.score(y_predicted, y_true)`
Calculate score from applying fitted objective to predicted values
If a higher score is better than a lower score, set `greater_is_better` attribute to True

Parameters

- **y_predicted** (*list*) – the predictions from the model. If `needs_proba` is True, it is the probability estimates
- **y_true** (*list*) – the ground truth for the predictions.

- **extra_cols** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if `uses_extra_columns` is `True`.

Returns score

`evalml.objectives.PrecisionMicro.supports_problem_type`

classmethod `PrecisionMicro.supports_problem_type` (*problem_type*)

Checks if objective supports given `ProblemType`

Parameters **problem_type** (*str* or *ProblemType*) – problem type to check

Returns whether objective supports `ProblemType`

Return type bool

`evalml.objectives.PrecisionMacro`

class `evalml.objectives.PrecisionMacro` (*verbose=False*)

Precision score for multiclass classification using macro averaging

Methods

<code>__init__</code>	Initialize self.
<code>fit</code>	Learn the objective function based on the predictions from a model.
<code>predict</code>	Apply the learned objective function to the output of a model.
<code>score</code>	Calculate score from applying fitted objective to predicted values
<code>supports_problem_type</code>	Checks if objective supports given <code>ProblemType</code>

`evalml.objectives.PrecisionMacro.__init__`

`PrecisionMacro.__init__` (*verbose=False*)

Initialize self. See `help(type(self))` for accurate signature.

`evalml.objectives.PrecisionMacro.fit`

`PrecisionMacro.fit` (*y_predicted, y_true, extra_cols=None*)

Learn the objective function based on the predictions from a model.

If `needs_fitting` is false, this method won't be called

Parameters

- **y_predicted** (*list*) – the predictions from the model. If `needs_proba` is `True`, it is the probability estimates
- **y_true** (*list*) – the ground truth for the predictions.
- **extra_cols** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if `uses_extra_columns` is `True`.

Returns self

evalml.objectives.PrecisionMacro.predict

PrecisionMacro.**predict** (*y_predicted*, *extra_cols=None*)

Apply the learned objective function to the output of a model.

If *needs_fitting* is false, this method won't be called

Parameters *y_predicted* – the prediction to transform to final prediction

Returns predictions

evalml.objectives.PrecisionMacro.score

PrecisionMacro.**score** (*y_predicted*, *y_true*)

Calculate score from applying fitted objective to predicted values

If a higher score is better than a lower score, set *greater_is_better* attribute to True

Parameters

- **y_predicted** (*list*) – the predictions from the model. If *needs_proba* is True, it is the probability estimates
- **y_true** (*list*) – the ground truth for the predictions.
- **extra_cols** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if *uses_extra_columns* is True.

Returns score

evalml.objectives.PrecisionMacro.supports_problem_type

classmethod PrecisionMacro.**supports_problem_type** (*problem_type*)

Checks if objective supports given ProblemType

Parameters *problem_type* (*str* or *ProblemType*) – problem type to check

Returns whether objective supports ProblemType

Return type bool

evalml.objectives.PrecisionWeighted

class evalml.objectives.PrecisionWeighted (*verbose=False*)

Precision score for multiclass classification using weighted averaging

Methods

<code>__init__</code>	Initialize self.
<code>fit</code>	Learn the objective function based on the predictions from a model.

Continued on next page

Table 42 – continued from previous page

<i>predict</i>	Apply the learned objective function to the output of a model.
<i>score</i>	Calculate score from applying fitted objective to predicted values
<i>supports_problem_type</i>	Checks if objective supports given ProblemType

evalml.objectives.PrecisionWeighted.__init__

PrecisionWeighted.**__init__**(*verbose=False*)
 Initialize self. See help(type(self)) for accurate signature.

evalml.objectives.PrecisionWeighted.fit

PrecisionWeighted.**fit**(*y_predicted, y_true, extra_cols=None*)
 Learn the objective function based on the predictions from a model.
 If needs_fitting is false, this method won't be called

Parameters

- **y_predicted** (*list*) – the predictions from the model. If needs_proba is True, it is the probability estimates
- **y_true** (*list*) – the ground truth for the predictions.
- **extra_cols** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if uses_extra_columns is True.

Returns self

evalml.objectives.PrecisionWeighted.predict

PrecisionWeighted.**predict**(*y_predicted, extra_cols=None*)
 Apply the learned objective function to the output of a model.
 If needs_fitting is false, this method won't be called

Parameters **y_predicted** – the prediction to transform to final prediction

Returns predictions

evalml.objectives.PrecisionWeighted.score

PrecisionWeighted.**score**(*y_predicted, y_true*)
 Calculate score from applying fitted objective to predicted values
 If a higher score is better than a lower score, set greater_is_better attribute to True

Parameters

- **y_predicted** (*list*) – the predictions from the model. If needs_proba is True, it is the probability estimates
- **y_true** (*list*) – the ground truth for the predictions.

- **extra_cols** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if `uses_extra_columns` is `True`.

Returns score

`evalml.objectives.PrecisionWeighted.supports_problem_type`

classmethod `PrecisionWeighted.supports_problem_type` (*problem_type*)

Checks if objective supports given `ProblemType`

Parameters **problem_type** (*str or ProblemType*) – problem type to check

Returns whether objective supports `ProblemType`

Return type bool

`evalml.objectives.Recall`

class `evalml.objectives.Recall` (*verbose=False*)

Recall score for binary classification

Methods

<code>__init__</code>	Initialize self.
<code>fit</code>	Learn the objective function based on the predictions from a model.
<code>predict</code>	Apply the learned objective function to the output of a model.
<code>score</code>	Calculate score from applying fitted objective to predicted values
<code>supports_problem_type</code>	Checks if objective supports given <code>ProblemType</code>

`evalml.objectives.Recall.__init__`

`Recall.__init__` (*verbose=False*)

Initialize self. See `help(type(self))` for accurate signature.

`evalml.objectives.Recall.fit`

`Recall.fit` (*y_predicted, y_true, extra_cols=None*)

Learn the objective function based on the predictions from a model.

If `needs_fitting` is false, this method won't be called

Parameters

- **y_predicted** (*list*) – the predictions from the model. If `needs_proba` is `True`, it is the probability estimates
- **y_true** (*list*) – the ground truth for the predictions.
- **extra_cols** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if `uses_extra_columns` is `True`.

Returns self

evalml.objectives.Recall.predict

`Recall.predict(y_predicted, extra_cols=None)`

Apply the learned objective function to the output of a model.

If `needs_fitting` is false, this method won't be called

Parameters `y_predicted` – the prediction to transform to final prediction

Returns predictions

evalml.objectives.Recall.score

`Recall.score(y_predicted, y_true)`

Calculate score from applying fitted objective to predicted values

If a higher score is better than a lower score, set `greater_is_better` attribute to True

Parameters

- `y_predicted` (*list*) – the predictions from the model. If `needs_proba` is True, it is the probability estimates
- `y_true` (*list*) – the ground truth for the predictions.
- `extra_cols` (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if `uses_extra_columns` is True.

Returns score

evalml.objectives.Recall.supports_problem_type

classmethod `Recall.supports_problem_type(problem_type)`

Checks if objective supports given `ProblemType`

Parameters `problem_type` (*str or ProblemType*) – problem type to check

Returns whether objective supports `ProblemType`

Return type bool

evalml.objectives.RecallMicro

class `evalml.objectives.RecallMicro(verbose=False)`

Recall score for multiclass classification using micro averaging

Methods

<code>__init__</code>	Initialize self.
<code>fit</code>	Learn the objective function based on the predictions from a model.

Continued on next page

Table 44 – continued from previous page

<i>predict</i>	Apply the learned objective function to the output of a model.
<i>score</i>	Calculate score from applying fitted objective to predicted values
<i>supports_problem_type</i>	Checks if objective supports given ProblemType

evalml.objectives.RecallMicro.__init__

`RecallMicro.__init__(verbose=False)`

Initialize self. See help(type(self)) for accurate signature.

evalml.objectives.RecallMicro.fit

`RecallMicro.fit(y_predicted, y_true, extra_cols=None)`

Learn the objective function based on the predictions from a model.

If needs_fitting is false, this method won't be called

Parameters

- **y_predicted** (*list*) – the predictions from the model. If needs_proba is True, it is the probability estimates
- **y_true** (*list*) – the ground truth for the predictions.
- **extra_cols** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if uses_extra_columns is True.

Returns self

evalml.objectives.RecallMicro.predict

`RecallMicro.predict(y_predicted, extra_cols=None)`

Apply the learned objective function to the output of a model.

If needs_fitting is false, this method won't be called

Parameters **y_predicted** – the prediction to transform to final prediction

Returns predictions

evalml.objectives.RecallMicro.score

`RecallMicro.score(y_predicted, y_true)`

Calculate score from applying fitted objective to predicted values

If a higher score is better than a lower score, set greater_is_better attribute to True

Parameters

- **y_predicted** (*list*) – the predictions from the model. If needs_proba is True, it is the probability estimates
- **y_true** (*list*) – the ground truth for the predictions.

- **extra_cols** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if `uses_extra_columns` is `True`.

Returns score

`evalml.objectives.RecallMicro.supports_problem_type`

classmethod `RecallMicro.supports_problem_type` (*problem_type*)

Checks if objective supports given `ProblemType`

Parameters **problem_type** (*str* or *ProblemType*) – problem type to check

Returns whether objective supports `ProblemType`

Return type bool

`evalml.objectives.RecallMacro`

class `evalml.objectives.RecallMacro` (*verbose=False*)

Recall score for multiclass classification using macro averaging

Methods

<code>__init__</code>	Initialize self.
<code>fit</code>	Learn the objective function based on the predictions from a model.
<code>predict</code>	Apply the learned objective function to the output of a model.
<code>score</code>	Calculate score from applying fitted objective to predicted values
<code>supports_problem_type</code>	Checks if objective supports given <code>ProblemType</code>

`evalml.objectives.RecallMacro.__init__`

`RecallMacro.__init__` (*verbose=False*)

Initialize self. See `help(type(self))` for accurate signature.

`evalml.objectives.RecallMacro.fit`

`RecallMacro.fit` (*y_predicted*, *y_true*, *extra_cols=None*)

Learn the objective function based on the predictions from a model.

If `needs_fitting` is false, this method won't be called

Parameters

- **y_predicted** (*list*) – the predictions from the model. If `needs_proba` is `True`, it is the probability estimates
- **y_true** (*list*) – the ground truth for the predictions.
- **extra_cols** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if `uses_extra_columns` is `True`.

Returns self

evalml.objectives.RecallMacro.predict

`RecallMacro.predict(y_predicted, extra_cols=None)`

Apply the learned objective function to the output of a model.

If `needs_fitting` is false, this method won't be called

Parameters `y_predicted` – the prediction to transform to final prediction

Returns predictions

evalml.objectives.RecallMacro.score

`RecallMacro.score(y_predicted, y_true)`

Calculate score from applying fitted objective to predicted values

If a higher score is better than a lower score, set `greater_is_better` attribute to True

Parameters

- `y_predicted` (*list*) – the predictions from the model. If `needs_proba` is True, it is the probability estimates
- `y_true` (*list*) – the ground truth for the predictions.
- `extra_cols` (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if `uses_extra_columns` is True.

Returns score

evalml.objectives.RecallMacro.supports_problem_type

classmethod `RecallMacro.supports_problem_type(problem_type)`

Checks if objective supports given `ProblemType`

Parameters `problem_type` (*str or ProblemType*) – problem type to check

Returns whether objective supports `ProblemType`

Return type bool

evalml.objectives.RecallWeighted

class `evalml.objectives.RecallWeighted(verbose=False)`

Recall score for multiclass classification using weighted averaging

Methods

<code>__init__</code>	Initialize self.
<code>fit</code>	Learn the objective function based on the predictions from a model.

Continued on next page

Table 46 – continued from previous page

<i>predict</i>	Apply the learned objective function to the output of a model.
<i>score</i>	Calculate score from applying fitted objective to predicted values
<i>supports_problem_type</i>	Checks if objective supports given ProblemType

evalml.objectives.RecallWeighted.__init__

RecallWeighted.**__init__**(*verbose=False*)
 Initialize self. See help(type(self)) for accurate signature.

evalml.objectives.RecallWeighted.fit

RecallWeighted.**fit**(*y_predicted, y_true, extra_cols=None*)
 Learn the objective function based on the predictions from a model.
 If needs_fitting is false, this method won't be called

Parameters

- **y_predicted** (*list*) – the predictions from the model. If needs_proba is True, it is the probability estimates
- **y_true** (*list*) – the ground truth for the predictions.
- **extra_cols** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if uses_extra_columns is True.

Returns self

evalml.objectives.RecallWeighted.predict

RecallWeighted.**predict**(*y_predicted, extra_cols=None*)
 Apply the learned objective function to the output of a model.
 If needs_fitting is false, this method won't be called

Parameters **y_predicted** – the prediction to transform to final prediction

Returns predictions

evalml.objectives.RecallWeighted.score

RecallWeighted.**score**(*y_predicted, y_true*)
 Calculate score from applying fitted objective to predicted values
 If a higher score is better than a lower score, set greater_is_better attribute to True

Parameters

- **y_predicted** (*list*) – the predictions from the model. If needs_proba is True, it is the probability estimates
- **y_true** (*list*) – the ground truth for the predictions.

- **extra_cols** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if `uses_extra_columns` is `True`.

Returns score

`evalml.objectives.RecallWeighted.supports_problem_type`

classmethod `RecallWeighted.supports_problem_type` (*problem_type*)

Checks if objective supports given `ProblemType`

Parameters **problem_type** (*str* or *ProblemType*) – problem type to check

Returns whether objective supports `ProblemType`

Return type bool

`evalml.objectives.AUC`

class `evalml.objectives.AUC` (*verbose=False*)

AUC score for binary classification

Methods

<code>__init__</code>	Initialize self.
<code>fit</code>	Learn the objective function based on the predictions from a model.
<code>predict</code>	Apply the learned objective function to the output of a model.
<code>score</code>	Calculate score from applying fitted objective to predicted values
<code>supports_problem_type</code>	Checks if objective supports given <code>ProblemType</code>

`evalml.objectives.AUC.__init__`

`AUC.__init__` (*verbose=False*)

Initialize self. See `help(type(self))` for accurate signature.

`evalml.objectives.AUC.fit`

`AUC.fit` (*y_predicted*, *y_true*, *extra_cols=None*)

Learn the objective function based on the predictions from a model.

If `needs_fitting` is false, this method won't be called

Parameters

- **y_predicted** (*list*) – the predictions from the model. If `needs_proba` is `True`, it is the probability estimates
- **y_true** (*list*) – the ground truth for the predictions.
- **extra_cols** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if `uses_extra_columns` is `True`.

Returns self

evalml.objectives.AUC.predict

AUC.predict (*y_predicted*, *extra_cols=None*)

Apply the learned objective function to the output of a model.

If `needs_fitting` is false, this method won't be called

Parameters *y_predicted* – the prediction to transform to final prediction

Returns predictions

evalml.objectives.AUC.score

AUC.score (*y_predicted*, *y_true*)

Calculate score from applying fitted objective to predicted values

If a higher score is better than a lower score, set `greater_is_better` attribute to True

Parameters

- **y_predicted** (*list*) – the predictions from the model. If `needs_proba` is True, it is the probability estimates
- **y_true** (*list*) – the ground truth for the predictions.
- **extra_cols** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if `uses_extra_columns` is True.

Returns score

evalml.objectives.AUC.supports_problem_type

classmethod **AUC.supports_problem_type** (*problem_type*)

Checks if objective supports given `ProblemType`

Parameters *problem_type* (*str* or *ProblemType*) – problem type to check

Returns whether objective supports `ProblemType`

Return type bool

evalml.objectives.AUCMicro

class **evalml.objectives.AUCMicro** (*verbose=False*)

AUC score for multiclass classification using micro averaging

Methods

<code>__init__</code>	Initialize self.
<code>fit</code>	Learn the objective function based on the predictions from a model.

Continued on next page

Table 48 – continued from previous page

<i>predict</i>	Apply the learned objective function to the output of a model.
<i>score</i>	Calculate score from applying fitted objective to predicted values
<i>supports_problem_type</i>	Checks if objective supports given ProblemType

evalml.objectives.AUCMicro.__init__

AUCMicro.**__init__**(*verbose=False*)

Initialize self. See help(type(self)) for accurate signature.

evalml.objectives.AUCMicro.fit

AUCMicro.**fit**(*y_predicted, y_true, extra_cols=None*)

Learn the objective function based on the predictions from a model.

If needs_fitting is false, this method won't be called

Parameters

- **y_predicted** (*list*) – the predictions from the model. If needs_proba is True, it is the probability estimates
- **y_true** (*list*) – the ground truth for the predictions.
- **extra_cols** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if uses_extra_columns is True.

Returns self

evalml.objectives.AUCMicro.predict

AUCMicro.**predict**(*y_predicted, extra_cols=None*)

Apply the learned objective function to the output of a model.

If needs_fitting is false, this method won't be called

Parameters **y_predicted** – the prediction to transform to final prediction

Returns predictions

evalml.objectives.AUCMicro.score

AUCMicro.**score**(*y_predicted, y_true*)

Calculate score from applying fitted objective to predicted values

If a higher score is better than a lower score, set greater_is_better attribute to True

Parameters

- **y_predicted** (*list*) – the predictions from the model. If needs_proba is True, it is the probability estimates
- **y_true** (*list*) – the ground truth for the predictions.

- **extra_cols** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if `uses_extra_columns` is `True`.

Returns score

`evalml.objectives.AUCMicro.supports_problem_type`

classmethod `AUCMicro.supports_problem_type` (*problem_type*)

Checks if objective supports given `ProblemType`

Parameters **problem_type** (*str* or *ProblemType*) – problem type to check

Returns whether objective supports `ProblemType`

Return type bool

`evalml.objectives.AUCMacro`

class `evalml.objectives.AUCMacro` (*verbose=False*)

AUC score for multiclass classification using macro averaging

Methods

<code>__init__</code>	Initialize self.
<code>fit</code>	Learn the objective function based on the predictions from a model.
<code>predict</code>	Apply the learned objective function to the output of a model.
<code>score</code>	Calculate score from applying fitted objective to predicted values
<code>supports_problem_type</code>	Checks if objective supports given <code>ProblemType</code>

`evalml.objectives.AUCMacro.__init__`

`AUCMacro.__init__` (*verbose=False*)

Initialize self. See `help(type(self))` for accurate signature.

`evalml.objectives.AUCMacro.fit`

`AUCMacro.fit` (*y_predicted, y_true, extra_cols=None*)

Learn the objective function based on the predictions from a model.

If `needs_fitting` is false, this method won't be called

Parameters

- **y_predicted** (*list*) – the predictions from the model. If `needs_proba` is `True`, it is the probability estimates
- **y_true** (*list*) – the ground truth for the predictions.
- **extra_cols** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if `uses_extra_columns` is `True`.

Returns self

evalml.objectives.AUCMacro.predict

`AUCMacro.predict(y_predicted, extra_cols=None)`

Apply the learned objective function to the output of a model.

If `needs_fitting` is false, this method won't be called

Parameters `y_predicted` – the prediction to transform to final prediction

Returns predictions

evalml.objectives.AUCMacro.score

`AUCMacro.score(y_predicted, y_true)`

Calculate score from applying fitted objective to predicted values

If a higher score is better than a lower score, set `greater_is_better` attribute to True

Parameters

- `y_predicted` (*list*) – the predictions from the model. If `needs_proba` is True, it is the probability estimates
- `y_true` (*list*) – the ground truth for the predictions.
- `extra_cols` (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if `uses_extra_columns` is True.

Returns score

evalml.objectives.AUCMacro.supports_problem_type

classmethod `AUCMacro.supports_problem_type(problem_type)`

Checks if objective supports given `ProblemType`

Parameters `problem_type` (*str or ProblemType*) – problem type to check

Returns whether objective supports `ProblemType`

Return type bool

evalml.objectives.AUCWeighted

class `evalml.objectives.AUCWeighted(verbose=False)`

AUC Score for multiclass classification using weighted averaging

Methods

<code>__init__</code>	Initialize self.
<code>fit</code>	Learn the objective function based on the predictions from a model.

Continued on next page

Table 50 – continued from previous page

<i>predict</i>	Apply the learned objective function to the output of a model.
<i>score</i>	Calculate score from applying fitted objective to predicted values
<i>supports_problem_type</i>	Checks if objective supports given ProblemType

evalml.objectives.AUCWeighted.__init__

AUCWeighted.__init__(*verbose=False*)

Initialize self. See help(type(self)) for accurate signature.

evalml.objectives.AUCWeighted.fit

AUCWeighted.fit(*y_predicted*, *y_true*, *extra_cols=None*)

Learn the objective function based on the predictions from a model.

If *needs_fitting* is false, this method won't be called

Parameters

- **y_predicted** (*list*) – the predictions from the model. If *needs_proba* is True, it is the probability estimates
- **y_true** (*list*) – the ground truth for the predictions.
- **extra_cols** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if *uses_extra_columns* is True.

Returns self

evalml.objectives.AUCWeighted.predict

AUCWeighted.predict(*y_predicted*, *extra_cols=None*)

Apply the learned objective function to the output of a model.

If *needs_fitting* is false, this method won't be called

Parameters **y_predicted** – the prediction to transform to final prediction

Returns predictions

evalml.objectives.AUCWeighted.score

AUCWeighted.score(*y_predicted*, *y_true*)

Calculate score from applying fitted objective to predicted values

If a higher score is better than a lower score, set *greater_is_better* attribute to True

Parameters

- **y_predicted** (*list*) – the predictions from the model. If *needs_proba* is True, it is the probability estimates
- **y_true** (*list*) – the ground truth for the predictions.

- **extra_cols** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if `uses_extra_columns` is `True`.

Returns score

`evalml.objectives.AUCWeighted.supports_problem_type`

classmethod `AUCWeighted.supports_problem_type(problem_type)`

Checks if objective supports given `ProblemType`

Parameters `problem_type` (*str* or *ProblemType*) – problem type to check

Returns whether objective supports `ProblemType`

Return type bool

`evalml.objectives.LogLoss`

class `evalml.objectives.LogLoss(verbose=False)`

Log Loss for both binary and multiclass classification

Methods

<code>__init__</code>	Initialize self.
<code>fit</code>	Learn the objective function based on the predictions from a model.
<code>predict</code>	Apply the learned objective function to the output of a model.
<code>score</code>	Calculate score from applying fitted objective to predicted values
<code>supports_problem_type</code>	Checks if objective supports given <code>ProblemType</code>

`evalml.objectives.LogLoss.__init__`

`LogLoss.__init__(verbose=False)`

Initialize self. See `help(type(self))` for accurate signature.

`evalml.objectives.LogLoss.fit`

`LogLoss.fit(y_predicted, y_true, extra_cols=None)`

Learn the objective function based on the predictions from a model.

If `needs_fitting` is false, this method won't be called

Parameters

- **y_predicted** (*list*) – the predictions from the model. If `needs_proba` is `True`, it is the probability estimates
- **y_true** (*list*) – the ground truth for the predictions.
- **extra_cols** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if `uses_extra_columns` is `True`.

Returns self

evalml.objectives.LogLoss.predict

`LogLoss.predict(y_predicted, extra_cols=None)`

Apply the learned objective function to the output of a model.

If `needs_fitting` is false, this method won't be called

Parameters `y_predicted` – the prediction to transform to final prediction

Returns predictions

evalml.objectives.LogLoss.score

`LogLoss.score(y_predicted, y_true)`

Calculate score from applying fitted objective to predicted values

If a higher score is better than a lower score, set `greater_is_better` attribute to True

Parameters

- `y_predicted` (*list*) – the predictions from the model. If `needs_proba` is True, it is the probability estimates
- `y_true` (*list*) – the ground truth for the predictions.
- `extra_cols` (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if `uses_extra_columns` is True.

Returns score

evalml.objectives.LogLoss.supports_problem_type

classmethod `LogLoss.supports_problem_type(problem_type)`

Checks if objective supports given `ProblemType`

Parameters `problem_type` (*str or ProblemType*) – problem type to check

Returns whether objective supports `ProblemType`

Return type bool

evalml.objectives.MCC

class `evalml.objectives.MCC(verbose=False)`

Matthews correlation coefficient for both binary and multiclass classification

Methods

<code>__init__</code>	Initialize self.
<code>fit</code>	Learn the objective function based on the predictions from a model.

Continued on next page

Table 52 – continued from previous page

<code>predict</code>	Apply the learned objective function to the output of a model.
<code>score</code>	Calculate score from applying fitted objective to predicted values
<code>supports_problem_type</code>	Checks if objective supports given ProblemType

`evalml.objectives.MCC.__init__`

`MCC.__init__(verbose=False)`

Initialize self. See `help(type(self))` for accurate signature.

`evalml.objectives.MCC.fit`

`MCC.fit(y_predicted, y_true, extra_cols=None)`

Learn the objective function based on the predictions from a model.

If `needs_fitting` is false, this method won't be called

Parameters

- **`y_predicted`** (*list*) – the predictions from the model. If `needs_proba` is True, it is the probability estimates
- **`y_true`** (*list*) – the ground truth for the predictions.
- **`extra_cols`** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if `uses_extra_columns` is True.

Returns self

`evalml.objectives.MCC.predict`

`MCC.predict(y_predicted, extra_cols=None)`

Apply the learned objective function to the output of a model.

If `needs_fitting` is false, this method won't be called

Parameters **`y_predicted`** – the prediction to transform to final prediction

Returns predictions

`evalml.objectives.MCC.score`

`MCC.score(y_predicted, y_true)`

Calculate score from applying fitted objective to predicted values

If a higher score is better than a lower score, set `greater_is_better` attribute to True

Parameters

- **`y_predicted`** (*list*) – the predictions from the model. If `needs_proba` is True, it is the probability estimates
- **`y_true`** (*list*) – the ground truth for the predictions.

- **extra_cols** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if `uses_extra_columns` is `True`.

Returns score

`evalml.objectives.MCC.supports_problem_type`

classmethod `MCC.supports_problem_type` (*problem_type*)

Checks if objective supports given `ProblemType`

Parameters **problem_type** (*str* or *ProblemType*) – problem type to check

Returns whether objective supports `ProblemType`

Return type bool

`evalml.objectives.ROC`

class `evalml.objectives.ROC` (*verbose=False*)

Receiver Operating Characteristic score for binary classification.

Methods

<code>__init__</code>	Initialize self.
<code>fit</code>	Learn the objective function based on the predictions from a model.
<code>predict</code>	Apply the learned objective function to the output of a model.
<code>score</code>	Calculate score from applying fitted objective to predicted values
<code>supports_problem_type</code>	Checks if objective supports given <code>ProblemType</code>

`evalml.objectives.ROC.__init__`

`ROC.__init__` (*verbose=False*)

Initialize self. See `help(type(self))` for accurate signature.

`evalml.objectives.ROC.fit`

`ROC.fit` (*y_predicted, y_true, extra_cols=None*)

Learn the objective function based on the predictions from a model.

If `needs_fitting` is false, this method won't be called

Parameters

- **y_predicted** (*list*) – the predictions from the model. If `needs_proba` is `True`, it is the probability estimates
- **y_true** (*list*) – the ground truth for the predictions.
- **extra_cols** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if `uses_extra_columns` is `True`.

Returns self

`evalml.objectives.ROC.predict`

`ROC.predict` (*y_predicted*, *extra_cols=None*)

Apply the learned objective function to the output of a model.

If `needs_fitting` is false, this method won't be called

Parameters `y_predicted` – the prediction to transform to final prediction

Returns predictions

`evalml.objectives.ROC.score`

`ROC.score` (*y_predicted*, *y_true*)

Calculate score from applying fitted objective to predicted values

If a higher score is better than a lower score, set `greater_is_better` attribute to True

Parameters

- `y_predicted` (*list*) – the predictions from the model. If `needs_proba` is True, it is the probability estimates
- `y_true` (*list*) – the ground truth for the predictions.
- `extra_cols` (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if `uses_extra_columns` is True.

Returns score

`evalml.objectives.ROC.supports_problem_type`

classmethod `ROC.supports_problem_type` (*problem_type*)

Checks if objective supports given `ProblemType`

Parameters `problem_type` (*str or ProblemType*) – problem type to check

Returns whether objective supports `ProblemType`

Return type bool

`evalml.objectives.ConfusionMatrix`

class `evalml.objectives.ConfusionMatrix` (*verbose=False*)

Confusion matrix for classification problems

Methods

<code>__init__</code>	Initialize self.
<code>fit</code>	Learn the objective function based on the predictions from a model.

Continued on next page

Table 54 – continued from previous page

<code>predict</code>	Apply the learned objective function to the output of a model.
<code>score</code>	Calculate score from applying fitted objective to predicted values
<code>supports_problem_type</code>	Checks if objective supports given ProblemType

evalml.objectives.ConfusionMatrix.__init__

`ConfusionMatrix.__init__(verbose=False)`

Initialize self. See help(type(self)) for accurate signature.

evalml.objectives.ConfusionMatrix.fit

`ConfusionMatrix.fit(y_predicted, y_true, extra_cols=None)`

Learn the objective function based on the predictions from a model.

If needs_fitting is false, this method won't be called

Parameters

- **y_predicted** (*list*) – the predictions from the model. If needs_proba is True, it is the probability estimates
- **y_true** (*list*) – the ground truth for the predictions.
- **extra_cols** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if uses_extra_columns is True.

Returns self

evalml.objectives.ConfusionMatrix.predict

`ConfusionMatrix.predict(y_predicted, extra_cols=None)`

Apply the learned objective function to the output of a model.

If needs_fitting is false, this method won't be called

Parameters **y_predicted** – the prediction to transform to final prediction

Returns predictions

evalml.objectives.ConfusionMatrix.score

`ConfusionMatrix.score(y_predicted, y_true)`

Calculate score from applying fitted objective to predicted values

If a higher score is better than a lower score, set greater_is_better attribute to True

Parameters

- **y_predicted** (*list*) – the predictions from the model. If needs_proba is True, it is the probability estimates
- **y_true** (*list*) – the ground truth for the predictions.

- **extra_cols** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if `uses_extra_columns` is `True`.

Returns score

`evalml.objectives.ConfusionMatrix.supports_problem_type`

classmethod `ConfusionMatrix.supports_problem_type` (*problem_type*)

Checks if objective supports given `ProblemType`

Parameters `problem_type` (*str* or *ProblemType*) – problem type to check

Returns whether objective supports `ProblemType`

Return type bool

Regression

<i>R2</i>	Coefficient of determination for regression
<i>MAE</i>	Mean absolute error for regression
<i>MSE</i>	Mean squared error for regression
<i>MSLE</i>	Mean squared log error for regression
<i>MedianAE</i>	Median absolute error for regression
<i>MaxError</i>	Maximum residual error for regression
<i>ExpVariance</i>	Explained variance score for regression

`evalml.objectives.R2`

class `evalml.objectives.R2` (*verbose=False*)

Coefficient of determination for regression

Methods

<code>__init__</code>	Initialize self.
<code>fit</code>	Learn the objective function based on the predictions from a model.
<code>predict</code>	Apply the learned objective function to the output of a model.
<code>score</code>	Calculate score from applying fitted objective to predicted values
<code>supports_problem_type</code>	Checks if objective supports given <code>ProblemType</code>

`evalml.objectives.R2.__init__`

`R2.__init__` (*verbose=False*)

Initialize self. See `help(type(self))` for accurate signature.

evalml.objectives.R2.fit**R2.fit** (*y_predicted*, *y_true*, *extra_cols=None*)

Learn the objective function based on the predictions from a model.

If *needs_fitting* is false, this method won't be called**Parameters**

- **y_predicted** (*list*) – the predictions from the model. If *needs_proba* is True, it is the probability estimates
- **y_true** (*list*) – the ground truth for the predictions.
- **extra_cols** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if *uses_extra_columns* is True.

Returns *self***evalml.objectives.R2.predict****R2.predict** (*y_predicted*, *extra_cols=None*)

Apply the learned objective function to the output of a model.

If *needs_fitting* is false, this method won't be called**Parameters** **y_predicted** – the prediction to transform to final prediction**Returns** predictions**evalml.objectives.R2.score****R2.score** (*y_predicted*, *y_true*)

Calculate score from applying fitted objective to predicted values

If a higher score is better than a lower score, set *greater_is_better* attribute to True**Parameters**

- **y_predicted** (*list*) – the predictions from the model. If *needs_proba* is True, it is the probability estimates
- **y_true** (*list*) – the ground truth for the predictions.
- **extra_cols** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if *uses_extra_columns* is True.

Returns score**evalml.objectives.R2.supports_problem_type****classmethod** **R2.supports_problem_type** (*problem_type*)Checks if objective supports given *ProblemType***Parameters** **problem_type** (*str* or *ProblemType*) – problem type to check**Returns** whether objective supports *ProblemType***Return type** bool

evalml.objectives.MAE

class evalml.objectives.**MAE** (*verbose=False*)
Mean absolute error for regression

Methods

<code>__init__</code>	Initialize self.
<code>fit</code>	Learn the objective function based on the predictions from a model.
<code>predict</code>	Apply the learned objective function to the output of a model.
<code>score</code>	Calculate score from applying fitted objective to predicted values
<code>supports_problem_type</code>	Checks if objective supports given ProblemType

evalml.objectives.MAE.__init__

MAE.__init__ (*verbose=False*)
Initialize self. See help(type(self)) for accurate signature.

evalml.objectives.MAE.fit

MAE.fit (*y_predicted, y_true, extra_cols=None*)
Learn the objective function based on the predictions from a model.
If needs_fitting is false, this method won't be called

Parameters

- **y_predicted** (*list*) – the predictions from the model. If needs_proba is True, it is the probability estimates
- **y_true** (*list*) – the ground truth for the predictions.
- **extra_cols** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if uses_extra_columns is True.

Returns self

evalml.objectives.MAE.predict

MAE.predict (*y_predicted, extra_cols=None*)
Apply the learned objective function to the output of a model.
If needs_fitting is false, this method won't be called

Parameters **y_predicted** – the prediction to transform to final prediction

Returns predictions

evalml.objectives.MAE.score**MAE.score** (*y_predicted*, *y_true*)

Calculate score from applying fitted objective to predicted values

If a higher score is better than a lower score, set `greater_is_better` attribute to `True`**Parameters**

- **y_predicted** (*list*) – the predictions from the model. If `needs_proba` is `True`, it is the probability estimates
- **y_true** (*list*) – the ground truth for the predictions.
- **extra_cols** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if `uses_extra_columns` is `True`.

Returns score**evalml.objectives.MAE.supports_problem_type****classmethod** **MAE.supports_problem_type** (*problem_type*)Checks if objective supports given `ProblemType`**Parameters** **problem_type** (*str or ProblemType*) – problem type to check**Returns** whether objective supports `ProblemType`**Return type** bool**evalml.objectives.MSE****class** **evalml.objectives.MSE** (*verbose=False*)

Mean squared error for regression

Methods

<code>__init__</code>	Initialize self.
<code>fit</code>	Learn the objective function based on the predictions from a model.
<code>predict</code>	Apply the learned objective function to the output of a model.
<code>score</code>	Calculate score from applying fitted objective to predicted values
<code>supports_problem_type</code>	Checks if objective supports given <code>ProblemType</code>

evalml.objectives.MSE.__init__**MSE.__init__** (*verbose=False*)Initialize self. See `help(type(self))` for accurate signature.

`evalml.objectives.MSE.fit`

`MSE.fit` (*y_predicted*, *y_true*, *extra_cols=None*)

Learn the objective function based on the predictions from a model.

If `needs_fitting` is false, this method won't be called

Parameters

- **`y_predicted`** (*list*) – the predictions from the model. If `needs_proba` is True, it is the probability estimates
- **`y_true`** (*list*) – the ground truth for the predictions.
- **`extra_cols`** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if `uses_extra_columns` is True.

Returns `self`

`evalml.objectives.MSE.predict`

`MSE.predict` (*y_predicted*, *extra_cols=None*)

Apply the learned objective function to the output of a model.

If `needs_fitting` is false, this method won't be called

Parameters **`y_predicted`** – the prediction to transform to final prediction

Returns predictions

`evalml.objectives.MSE.score`

`MSE.score` (*y_predicted*, *y_true*)

Calculate score from applying fitted objective to predicted values

If a higher score is better than a lower score, set `greater_is_better` attribute to True

Parameters

- **`y_predicted`** (*list*) – the predictions from the model. If `needs_proba` is True, it is the probability estimates
- **`y_true`** (*list*) – the ground truth for the predictions.
- **`extra_cols`** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if `uses_extra_columns` is True.

Returns score

`evalml.objectives.MSE.supports_problem_type`

classmethod `MSE.supports_problem_type` (*problem_type*)

Checks if objective supports given `ProblemType`

Parameters **`problem_type`** (*str or ProblemType*) – problem type to check

Returns whether objective supports `ProblemType`

Return type bool

evalml.objectives.MSLE

class evalml.objectives.**MSLE** (*verbose=False*)
Mean squared log error for regression

Methods

<code>__init__</code>	Initialize self.
<code>fit</code>	Learn the objective function based on the predictions from a model.
<code>predict</code>	Apply the learned objective function to the output of a model.
<code>score</code>	Calculate score from applying fitted objective to predicted values
<code>supports_problem_type</code>	Checks if objective supports given ProblemType

evalml.objectives.MSLE.__init__

MSLE.__init__ (*verbose=False*)
Initialize self. See help(type(self)) for accurate signature.

evalml.objectives.MSLE.fit

MSLE.fit (*y_predicted, y_true, extra_cols=None*)
Learn the objective function based on the predictions from a model.
If needs_fitting is false, this method won't be called

Parameters

- **y_predicted** (*list*) – the predictions from the model. If needs_proba is True, it is the probability estimates
- **y_true** (*list*) – the ground truth for the predictions.
- **extra_cols** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if uses_extra_columns is True.

Returns self

evalml.objectives.MSLE.predict

MSLE.predict (*y_predicted, extra_cols=None*)
Apply the learned objective function to the output of a model.
If needs_fitting is false, this method won't be called

Parameters **y_predicted** – the prediction to transform to final prediction

Returns predictions

evalml.objectives.MSLE.score**MSLE.score** (*y_predicted*, *y_true*)

Calculate score from applying fitted objective to predicted values

If a higher score is better than a lower score, set `greater_is_better` attribute to `True`**Parameters**

- **y_predicted** (*list*) – the predictions from the model. If `needs_proba` is `True`, it is the probability estimates
- **y_true** (*list*) – the ground truth for the predictions.
- **extra_cols** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if `uses_extra_columns` is `True`.

Returns score**evalml.objectives.MSLE.supports_problem_type****classmethod** **MSLE.supports_problem_type** (*problem_type*)Checks if objective supports given `ProblemType`**Parameters** **problem_type** (*str or ProblemType*) – problem type to check**Returns** whether objective supports `ProblemType`**Return type** bool**evalml.objectives.MedianAE****class** **evalml.objectives.MedianAE** (*verbose=False*)

Median absolute error for regression

Methods

<code>__init__</code>	Initialize self.
<code>fit</code>	Learn the objective function based on the predictions from a model.
<code>predict</code>	Apply the learned objective function to the output of a model.
<code>score</code>	Calculate score from applying fitted objective to predicted values
<code>supports_problem_type</code>	Checks if objective supports given <code>ProblemType</code>

evalml.objectives.MedianAE.__init__**MedianAE.__init__** (*verbose=False*)Initialize self. See `help(type(self))` for accurate signature.

evalml.objectives.MedianAE.fit`MedianAE.fit(y_predicted, y_true, extra_cols=None)`

Learn the objective function based on the predictions from a model.

If needs_fitting is false, this method won't be called

Parameters

- **y_predicted** (*list*) – the predictions from the model. If needs_proba is True, it is the probability estimates
- **y_true** (*list*) – the ground truth for the predictions.
- **extra_cols** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if uses_extra_columns is True.

Returns self

evalml.objectives.MedianAE.predict`MedianAE.predict(y_predicted, extra_cols=None)`

Apply the learned objective function to the output of a model.

If needs_fitting is false, this method won't be called

Parameters **y_predicted** – the prediction to transform to final prediction

Returns predictions

evalml.objectives.MedianAE.score`MedianAE.score(y_predicted, y_true)`

Calculate score from applying fitted objective to predicted values

If a higher score is better than a lower score, set greater_is_better attribute to True

Parameters

- **y_predicted** (*list*) – the predictions from the model. If needs_proba is True, it is the probability estimates
- **y_true** (*list*) – the ground truth for the predictions.
- **extra_cols** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if uses_extra_columns is True.

Returns score

evalml.objectives.MedianAE.supports_problem_type`classmethod MedianAE.supports_problem_type(problem_type)`

Checks if objective supports given ProblemType

Parameters **problem_type** (*str or ProblemType*) – problem type to check

Returns whether objective supports ProblemType

Return type bool

evalml.objectives.MaxError

class evalml.objectives.**MaxError** (*verbose=False*)
Maximum residual error for regression

Methods

<code>__init__</code>	Initialize self.
<code>fit</code>	Learn the objective function based on the predictions from a model.
<code>predict</code>	Apply the learned objective function to the output of a model.
<code>score</code>	Calculate score from applying fitted objective to predicted values
<code>supports_problem_type</code>	Checks if objective supports given ProblemType

evalml.objectives.MaxError.__init__

`MaxError.__init__` (*verbose=False*)
Initialize self. See help(type(self)) for accurate signature.

evalml.objectives.MaxError.fit

`MaxError.fit` (*y_predicted, y_true, extra_cols=None*)
Learn the objective function based on the predictions from a model.
If needs_fitting is false, this method won't be called

Parameters

- **y_predicted** (*list*) – the predictions from the model. If needs_proba is True, it is the probability estimates
- **y_true** (*list*) – the ground truth for the predictions.
- **extra_cols** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if uses_extra_columns is True.

Returns self

evalml.objectives.MaxError.predict

`MaxError.predict` (*y_predicted, extra_cols=None*)
Apply the learned objective function to the output of a model.
If needs_fitting is false, this method won't be called

Parameters **y_predicted** – the prediction to transform to final prediction

Returns predictions

evalml.objectives.MaxError.score**MaxError.score** (*y_predicted*, *y_true*)

Calculate score from applying fitted objective to predicted values

If a higher score is better than a lower score, set `greater_is_better` attribute to `True`**Parameters**

- **y_predicted** (*list*) – the predictions from the model. If `needs_proba` is `True`, it is the probability estimates
- **y_true** (*list*) – the ground truth for the predictions.
- **extra_cols** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if `uses_extra_columns` is `True`.

Returns score**evalml.objectives.MaxError.supports_problem_type****classmethod** **MaxError.supports_problem_type** (*problem_type*)Checks if objective supports given `ProblemType`**Parameters** **problem_type** (*str or ProblemType*) – problem type to check**Returns** whether objective supports `ProblemType`**Return type** bool**evalml.objectives.ExpVariance****class** **evalml.objectives.ExpVariance** (*verbose=False*)

Explained variance score for regression

Methods

<code>__init__</code>	Initialize self.
<code>fit</code>	Learn the objective function based on the predictions from a model.
<code>predict</code>	Apply the learned objective function to the output of a model.
<code>score</code>	Calculate score from applying fitted objective to predicted values
<code>supports_problem_type</code>	Checks if objective supports given <code>ProblemType</code>

evalml.objectives.ExpVariance.__init__**ExpVariance.__init__** (*verbose=False*)Initialize self. See `help(type(self))` for accurate signature.

`evalml.objectives.ExpVariance.fit`

`ExpVariance.fit` (*y_predicted*, *y_true*, *extra_cols=None*)

Learn the objective function based on the predictions from a model.

If `needs_fitting` is false, this method won't be called

Parameters

- **`y_predicted`** (*list*) – the predictions from the model. If `needs_proba` is True, it is the probability estimates
- **`y_true`** (*list*) – the ground truth for the predictions.
- **`extra_cols`** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if `uses_extra_columns` is True.

Returns `self`

`evalml.objectives.ExpVariance.predict`

`ExpVariance.predict` (*y_predicted*, *extra_cols=None*)

Apply the learned objective function to the output of a model.

If `needs_fitting` is false, this method won't be called

Parameters **`y_predicted`** – the prediction to transform to final prediction

Returns predictions

`evalml.objectives.ExpVariance.score`

`ExpVariance.score` (*y_predicted*, *y_true*)

Calculate score from applying fitted objective to predicted values

If a higher score is better than a lower score, set `greater_is_better` attribute to True

Parameters

- **`y_predicted`** (*list*) – the predictions from the model. If `needs_proba` is True, it is the probability estimates
- **`y_true`** (*list*) – the ground truth for the predictions.
- **`extra_cols`** (*pd.DataFrame*) – any extra columns that are needed from training data to fit. Only provided if `uses_extra_columns` is True.

Returns score

`evalml.objectives.ExpVariance.supports_problem_type`

classmethod `ExpVariance.supports_problem_type` (*problem_type*)

Checks if objective supports given `ProblemType`

Parameters **`problem_type`** (*str or ProblemType*) – problem type to check

Returns whether objective supports `ProblemType`

Return type bool

Problem Types

<i>ProblemTypes</i>	Enum for type of machine learning problem: BINARY, MULTICLASS, or REGRESSION
---------------------	--

evalml.problem_types.ProblemTypes

class evalml.problem_types.**ProblemTypes**
Enum for type of machine learning problem: BINARY, MULTICLASS, or REGRESSION

<i>handle_problem_types</i>	Handles problem_type by either returning the ProblemTypes or converting from a str
-----------------------------	--

evalml.problem_types.handle_problem_types

evalml.problem_types.**handle_problem_types**(*problem_type*)
Handles problem_type by either returning the ProblemTypes or converting from a str

Parameters **problem_types** (*str* or *ProblemTypes*) – problem type that needs to be handled

Returns ProblemTypes

Tuners

<i>SKOptTuner</i>	Bayesian Optimizer
-------------------	--------------------

evalml.tuners.SKOptTuner

class evalml.tuners.**SKOptTuner**(*space*, *random_state=0*)
Bayesian Optimizer

Methods

<i>__init__</i>	Init SKOptTuner
<i>add</i>	Add score to sample
<i>propose</i>	Returns hyperparameters based off search space and samples

evalml.tuners.SKOptTuner.__init__

SKOptTuner.**__init__**(*space*, *random_state=0*)
Init SKOptTuner

Parameters

- **space** (*dict*) – search space for hyperparameters
- **random_state** (*int*) – random state

Returns self

Return type SKOptTuner

evalml.tuners.SKOptTuner.add

SKOptTuner.add(parameters, score)

Add score to sample

Parameters

- **parameters** (*dict*) – hyperparameters
- **score** (*float*) – associated score

Returns None

evalml.tuners.SKOptTuner.propose

SKOptTuner.propose()

Returns hyperparameters based off search space and samples

Returns proposed hyperparameters

Return type dict

Guardrails

<code>detect_highly_null</code>	Checks if there are any highly-null columns in a dataframe.
<code>detect_label_leakage</code>	Check if any of the features are highly correlated with the target.
<code>detect_outliers</code>	Checks if there are any outliers in a dataframe by using first Isolation Forest to obtain the anomaly score of each index and then using IQR to determine score anomalies.
<code>detect_id_columns</code>	Check if any of the features are ID columns.

evalml.guardrails.detect_highly_null

evalml.guardrails.detect_highly_null(X, percent_threshold=0.95)

Checks if there are any highly-null columns in a dataframe.

Parameters

- **X** (*pd.DataFrame*) – features
- **percent_threshold** (*float*) – Require that percentage of null values to be considered “highly-null”, defaults to .95

Returns A dictionary of features with column name or index and their percentage of null values

Example

```
>>> df = pd.DataFrame({
...     'lots_of_null': [None, None, None, None, 5],
...     'no_null': [1, 2, 3, 4, 5]
... })
>>> detect_highly_null(df, percent_threshold=0.8)
{'lots_of_null': 0.8}
```

evalml.guardrails.detect_label_leakage

evalml.guardrails.**detect_label_leakage**(X, y, threshold=0.95)

Check if any of the features are highly correlated with the target.

Currently only supports binary and numeric targets and features

Parameters

- **X** (*pd.DataFrame*) – The input features to check
- **y** (*pd.Series*) – the labels
- **threshold** (*float*) – the correlation threshold to be considered leakage. Defaults to .95

Returns leakage, dictionary of features with leakage and corresponding threshold

Example

```
>>> X = pd.DataFrame({
...     'leak': [10, 42, 31, 51, 61],
...     'x': [42, 54, 12, 64, 12],
...     'y': [12, 5, 13, 74, 24],
... })
>>> y = pd.Series([10, 42, 31, 51, 40])
>>> detect_label_leakage(X, y, threshold=0.8)
{'leak': 0.8827072320669518}
```

evalml.guardrails.detect_outliers

evalml.guardrails.**detect_outliers**(X, random_state=0)

Checks if there are any outliers in a dataframe by using first Isolation Forest to obtain the anomaly score of each index and then using IQR to determine score anomalies. Indices with score anomalies are considered outliers.

Parameters **X** (*pd.DataFrame*) – features

Returns A set of indices that may have outlier data.

Example

```
>>> df = pd.DataFrame({
...     'x': [1, 2, 3, 40, 5],
...     'y': [6, 7, 8, 990, 10],
...     'z': [-1, -2, -3, -1201, -4]
```

(continues on next page)

(continued from previous page)

```
... })
>>> detect_outliers(df)
[3]
```

evalml.guardrails.detect_id_columns

evalml.guardrails.**detect_id_columns**(X, threshold=1.0)

Check if any of the features are ID columns. Currently performs these simple checks:

- column name is “id”
- column name ends in “_id”
- column contains all unique values (and is not float / boolean)

Parameters

- **X** (*pd.DataFrame*) – The input features to check
- **threshold** (*float*) – the probability threshold to be considered an ID column. Defaults to 1.0

Returns A dictionary of features with column name or index and their probability of being ID columns

Example

```
>>> df = pd.DataFrame({
...     'df_id': [0, 1, 2, 3, 4],
...     'x': [10, 42, 31, 51, 61],
...     'y': [42, 54, 12, 64, 12]
... })
>>> detect_id_columns(df)
{'df_id': 1.0}
```

2.6.15 FAQ

What is the difference between EvalML and other AutoML libraries?

EvalML optimizes machine learning pipelines on *custom practical objectives* instead of vague machine learning loss functions so that it will find the best pipelines for your specific needs. Furthermore, EvalML *pipelines* are able to take in all kinds of data (missing values, categorical, etc.) as long as the data are in a single table. EvalML also allows you to build your own pipelines with existing or custom components so you can have more control over the AutoML process. Moreover, EvalML also provides you with support in the form of *guardrails* to ensure that you are aware of potential issues your data may cause with machine learning algorithms”.

How does EvalML handle missing values?

EvalML contains imputation components in its pipelines so that missing values are taken care of. EvalML optimizes over different types of imputation to search for the best possible pipeline. You can find more information about components [here](#) and in the API reference [here](#).

How does EvalML handle categorical encoding?

EvalML provides a *one-hot-encoding component* in its pipelines for categorical variables. EvalML plans to support other encoders in the future.

How does EvalML handle feature selection?

EvalML currently utilizes scikit-learn's `SelectFromModel` with a Random Forest classifier/regressor to handle feature selection. EvalML plans on supporting more feature selectors in the future. You can find more information in the API reference [here](#).

How are feature importances calculated?

Feature importance depends on the estimator used. Variable coefficients are used for regression-based estimators (Logistic Regression and Linear Regression) and Gini importance is used for tree-based estimators (Random Forest and XGBoost).

How does hyperparameter tuning work?

EvalML tunes hyperparameters for its pipelines through Bayesian optimization. In the future we plan to support more optimization techniques such as random search.

Can I create my own objective metric?

Yes you can! You can *create your own custom objective* so that EvalML optimizes the best model for your needs.

How does EvalML avoid overfitting?

EvalML provides *guardrails* to combat overfitting. Such guardrails include detecting label leakage, unstable pipelines, hold-out datasets and cross validation. EvalML defaults to using Stratified K-Fold cross-validation for classification problems and K-Fold cross-validation for regression problems but allows you to utilize your own cross-validation methods as well.

Can I create my own pipeline for EvalML?

Yes! EvalML allows you to create *custom pipelines* using modular components. This allows you to customize EvalML pipelines for your own needs or for AutoML.

Does EvalML work with X algorithm?

EvalML is constantly improving and adding new components and will allow your own algorithms to be used as components in our pipelines.

Symbols

`__init__()` (*evalml.AutoClassificationSearch method*), 54
`__init__()` (*evalml.AutoRegressionSearch method*), 55
`__init__()` (*evalml.objectives.AUC method*), 110
`__init__()` (*evalml.objectives.AUCMacro method*), 113
`__init__()` (*evalml.objectives.AUCMicro method*), 112
`__init__()` (*evalml.objectives.AUCWeighted method*), 115
`__init__()` (*evalml.objectives.ConfusionMatrix method*), 121
`__init__()` (*evalml.objectives.ExpVariance method*), 131
`__init__()` (*evalml.objectives.F1 method*), 92
`__init__()` (*evalml.objectives.F1Macro method*), 95
`__init__()` (*evalml.objectives.F1Micro method*), 94
`__init__()` (*evalml.objectives.F1Weighted method*), 97
`__init__()` (*evalml.objectives.FraudCost method*), 87
`__init__()` (*evalml.objectives.LeadScoring method*), 90
`__init__()` (*evalml.objectives.LogLoss method*), 116
`__init__()` (*evalml.objectives.MAE method*), 124
`__init__()` (*evalml.objectives.MCC method*), 118
`__init__()` (*evalml.objectives.MSE method*), 125
`__init__()` (*evalml.objectives.MSLE method*), 127
`__init__()` (*evalml.objectives.MaxError method*), 130
`__init__()` (*evalml.objectives.MedianAE method*), 128
`__init__()` (*evalml.objectives.Precision method*), 98
`__init__()` (*evalml.objectives.PrecisionMacro method*), 101
`__init__()` (*evalml.objectives.PrecisionMicro method*), 100
`__init__()` (*evalml.objectives.PrecisionWeighted method*), 103
`__init__()` (*evalml.objectives.R2 method*), 122
`__init__()` (*evalml.objectives.ROC method*), 119
`__init__()` (*evalml.objectives.Recall method*), 104
`__init__()` (*evalml.objectives.RecallMacro method*), 107
`__init__()` (*evalml.objectives.RecallMicro method*), 106
`__init__()` (*evalml.objectives.RecallWeighted method*), 109
`__init__()` (*evalml.pipelines.LinearRegressionPipeline method*), 85
`__init__()` (*evalml.pipelines.LogisticRegressionPipeline method*), 81
`__init__()` (*evalml.pipelines.PipelineBase method*), 74
`__init__()` (*evalml.pipelines.RFClassificationPipeline method*), 76
`__init__()` (*evalml.pipelines.RFRegressionPipeline method*), 83
`__init__()` (*evalml.pipelines.XGBoostPipeline method*), 78
`__init__()` (*evalml.pipelines.components.LinearRegressor method*), 71
`__init__()` (*evalml.pipelines.components.LogisticRegressionClassifier method*), 67
`__init__()` (*evalml.pipelines.components.OneHotEncoder method*), 58
`__init__()` (*evalml.pipelines.components.RFClassifierSelectFromModel method*), 62
`__init__()` (*evalml.pipelines.components.RFRegressorSelectFromModel method*), 60
`__init__()` (*evalml.pipelines.components.RandomForestClassifier method*), 68
`__init__()` (*evalml.pipelines.components.RandomForestRegressor method*), 72
`__init__()` (*evalml.pipelines.components.SimpleImputer method*), 64
`__init__()` (*evalml.pipelines.components.StandardScaler method*), 65
`__init__()` (*evalml.pipelines.components.XGBoostClassifier method*), 69
`__init__()` (*evalml.tuners.SKOptTuner method*), 133

A

`add()` (*evalml.tuners.SKOptTuner method*), 134
`AUC` (*class in evalml.objectives*), 110
`AUCMacro` (*class in evalml.objectives*), 113
`AUCMicro` (*class in evalml.objectives*), 111
`AUCWeighted` (*class in evalml.objectives*), 114
`AutoClassificationSearch` (*class in evalml*), 53
`AutoRegressionSearch` (*class in evalml*), 55

C

`ConfusionMatrix` (*class in evalml.objectives*), 120

D

`decision_function()`
 (*evalml.objectives.FraudCost method*), 88
`decision_function()`
 (*evalml.objectives.LeadScoring method*), 90
`describe()` (*evalml.pipelines.components.LinearRegressionClassifier method*), 71
`describe()` (*evalml.pipelines.components.LogisticRegressionClassifier method*), 67
`describe()` (*evalml.pipelines.components.OneHotEncoder method*), 58
`describe()` (*evalml.pipelines.components.RandomForestClassifier method*), 68
`describe()` (*evalml.pipelines.components.RandomForestRegressor method*), 72
`describe()` (*evalml.pipelines.components.RFClassifierSelector method*), 62
`describe()` (*evalml.pipelines.components.RFRegressorSelector method*), 60
`describe()` (*evalml.pipelines.components.SimpleImputer method*), 64
`describe()` (*evalml.pipelines.components.StandardScaler method*), 65
`describe()` (*evalml.pipelines.components.XGBoostClassifier method*), 69
`describe()` (*evalml.pipelines.LinearRegressionPipeline method*), 85
`describe()` (*evalml.pipelines.LogisticRegressionPipeline method*), 81
`describe()` (*evalml.pipelines.PipelineBase method*), 75
`describe()` (*evalml.pipelines.RFClassificationPipeline method*), 77
`describe()` (*evalml.pipelines.RFRegressionPipeline method*), 83
`describe()` (*evalml.pipelines.XGBoostPipeline method*), 79
`detect_highly_null()` (in module *evalml.guardrails*), 134
`detect_id_columns()` (in module *evalml.guardrails*), 136

`detect_label_leakage()` (in module *evalml.guardrails*), 135
`detect_outliers()` (in module *evalml.guardrails*), 135

E

`ExpVariance` (*class in evalml.objectives*), 131

F

`F1` (*class in evalml.objectives*), 92
`F1Macro` (*class in evalml.objectives*), 95
`F1Micro` (*class in evalml.objectives*), 93
`F1Weighted` (*class in evalml.objectives*), 96
`fit()` (*evalml.objectives.AUC method*), 110
`fit()` (*evalml.objectives.AUCMacro method*), 113
`fit()` (*evalml.objectives.AUCMicro method*), 112
`fit()` (*evalml.objectives.AUCWeighted method*), 115
`fit()` (*evalml.objectives.ConfusionMatrix method*), 121
`fit()` (*evalml.objectives.ExpVariance method*), 132
`fit()` (*evalml.objectives.F1 method*), 92
`fit()` (*evalml.objectives.F1Macro method*), 95
`fit()` (*evalml.objectives.F1Micro method*), 94
`fit()` (*evalml.objectives.F1Weighted method*), 97
`fit()` (*evalml.objectives.FraudCost method*), 88
`fit()` (*evalml.objectives.LeadScoring method*), 90
`fit()` (*evalml.objectives.LogLoss method*), 116
`fit()` (*evalml.objectives.MAE method*), 124
`fit()` (*evalml.objectives.MaxError method*), 130
`fit()` (*evalml.objectives.MCC method*), 118
`fit()` (*evalml.objectives.MedianAE method*), 129
`fit()` (*evalml.objectives.MSE method*), 126
`fit()` (*evalml.objectives.MSLE method*), 127
`fit()` (*evalml.objectives.Precision method*), 98
`fit()` (*evalml.objectives.PrecisionMacro method*), 101
`fit()` (*evalml.objectives.PrecisionMicro method*), 100
`fit()` (*evalml.objectives.PrecisionWeighted method*), 103
`fit()` (*evalml.objectives.R2 method*), 123
`fit()` (*evalml.objectives.Recall method*), 104
`fit()` (*evalml.objectives.RecallMacro method*), 107
`fit()` (*evalml.objectives.RecallMicro method*), 106
`fit()` (*evalml.objectives.RecallWeighted method*), 109
`fit()` (*evalml.objectives.ROC method*), 119
`fit()` (*evalml.pipelines.components.LinearRegressor method*), 71
`fit()` (*evalml.pipelines.components.LogisticRegressionClassifier method*), 67
`fit()` (*evalml.pipelines.components.OneHotEncoder method*), 59
`fit()` (*evalml.pipelines.components.RandomForestClassifier method*), 68
`fit()` (*evalml.pipelines.components.RandomForestRegressor method*), 72

`fit()` (`evalml.pipelines.components.RFClassifierSelectFromModel` method), 62
`fit()` (`evalml.pipelines.components.RFRegressorSelectFromModel` method), 60
`fit()` (`evalml.pipelines.components.SimpleImputer` method), 64
`fit()` (`evalml.pipelines.components.StandardScaler` method), 65
`fit()` (`evalml.pipelines.components.XGBoostClassifier` method), 70
`fit()` (`evalml.pipelines.LinearRegressionPipeline` method), 85
`fit()` (`evalml.pipelines.LogisticRegressionPipeline` method), 81
`fit()` (`evalml.pipelines.PipelineBase` method), 75
`fit()` (`evalml.pipelines.RFClassificationPipeline` method), 77
`fit()` (`evalml.pipelines.RFRegressionPipeline` method), 83
`fit()` (`evalml.pipelines.XGBoostPipeline` method), 79
`fit_transform()` (`evalml.pipelines.components.OneHotEncoder` method), 59
`fit_transform()` (`evalml.pipelines.components.RFClassifierSelectFromModel` method), 63
`fit_transform()` (`evalml.pipelines.components.RFRegressorSelectFromModel` method), 61
`fit_transform()` (`evalml.pipelines.components.SimpleImputer` method), 64
`fit_transform()` (`evalml.pipelines.components.StandardScaler` method), 66
FraudCost (class in `evalml.objectives`), 87

G

`generate_confusion_matrix()` (`evalml.AutoClassificationSearch.plot` method), 57
`generate_confusion_matrix()` (`evalml.AutoRegressionSearch.plot` method), 57
`generate_roc_plot()` (`evalml.AutoClassificationSearch.plot` method), 56
`generate_roc_plot()` (`evalml.AutoRegressionSearch.plot` method), 57
`get_component()` (`evalml.pipelines.LinearRegressionPipeline` method), 86
`get_component()` (`evalml.pipelines.LogisticRegressionPipeline` method), 81
`get_component()` (`evalml.pipelines.PipelineBase` method), 75
`get_component()` (`evalml.pipelines.RFClassificationPipeline` method), 77

`get_component()` (`evalml.pipelines.RFRegressionPipeline` method), 83
`get_confusion_matrix_data()` (`evalml.AutoClassificationSearch.plot` method), 57
`get_confusion_matrix_data()` (`evalml.AutoRegressionSearch.plot` method), 57
`get_feature_names()` (`evalml.pipelines.components.OneHotEncoder` method), 59
`get_indices()` (`evalml.pipelines.components.RFClassifierSelectFromModel` method), 63
`get_indices()` (`evalml.pipelines.components.RFRegressorSelectFromModel` method), 61
`get_names()` (`evalml.pipelines.components.RFClassifierSelectFromModel` method), 63
`get_names()` (`evalml.pipelines.components.RFRegressorSelectFromModel` method), 61
`get_pipelines()` (in module `evalml.pipelines`), 73
`select_from_model()` (`evalml.AutoClassificationSearch.plot` method), 56
`select_from_model()` (`evalml.AutoRegressionSearch.plot` method), 57

H

`load_data()` (in module `evalml.problem_types`), 133

L

LeadScoring (class in `evalml.objectives`), 89
LinearRegressionPipeline (class in `evalml.pipelines`), 84
LinearRegressor (class in `evalml.pipelines.components`), 70
`list_model_types()` (in module `evalml`), 58
`load_breast_cancer()` (in module `evalml.demos`), 52
`load_data()` (in module `evalml.preprocessing`), 52
`load_diabetes()` (in module `evalml.demos`), 52
`load_fraud()` (in module `evalml.demos`), 52
`load_pipeline()` (in module `evalml.pipelines`), 74
`load_wine()` (in module `evalml.demos`), 52

LogisticRegressionClassifier (class in `evalml.pipelines.components`), 66
LogisticRegressionPipeline (class in `evalml.pipelines`), 80

LogLoss (class in `evalml.objectives`), 116

M

MAE (class in `evalml.objectives`), 124
MaxError (class in `evalml.objectives`), 130

MCC (*class in evalml.objectives*), 117
 MedianAE (*class in evalml.objectives*), 128
 MSE (*class in evalml.objectives*), 125
 MSLE (*class in evalml.objectives*), 127

O

`objective_function()`
 (*evalml.objectives.FraudCost method*), 88
`objective_function()`
 (*evalml.objectives.LeadScoring method*), 90
 OneHotEncoder (*class in evalml.pipelines.components*), 58

P

PipelineBase (*class in evalml.pipelines*), 74
`plot()` (*evalml.pipelines.PipelineBase method*), 87
 Precision (*class in evalml.objectives*), 98
 PrecisionMacro (*class in evalml.objectives*), 101
 PrecisionMicro (*class in evalml.objectives*), 99
 PrecisionWeighted (*class in evalml.objectives*), 102
`predict()` (*evalml.objectives.AUC method*), 111
`predict()` (*evalml.objectives.AUCMacro method*), 114
`predict()` (*evalml.objectives.AUCMicro method*), 112
`predict()` (*evalml.objectives.AUCWeighted method*), 115
`predict()` (*evalml.objectives.ConfusionMatrix method*), 121
`predict()` (*evalml.objectives.ExpVariance method*), 132
`predict()` (*evalml.objectives.F1 method*), 93
`predict()` (*evalml.objectives.F1Macro method*), 96
`predict()` (*evalml.objectives.F1Micro method*), 94
`predict()` (*evalml.objectives.F1Weighted method*), 97
`predict()` (*evalml.objectives.FraudCost method*), 89
`predict()` (*evalml.objectives.LeadScoring method*), 91
`predict()` (*evalml.objectives.LogLoss method*), 117
`predict()` (*evalml.objectives.MAE method*), 124
`predict()` (*evalml.objectives.MaxError method*), 130
`predict()` (*evalml.objectives.MCC method*), 118
`predict()` (*evalml.objectives.MedianAE method*), 129
`predict()` (*evalml.objectives.MSE method*), 126
`predict()` (*evalml.objectives.MSLE method*), 127
`predict()` (*evalml.objectives.Precision method*), 99
`predict()` (*evalml.objectives.PrecisionMacro method*), 102
`predict()` (*evalml.objectives.PrecisionMicro method*), 100
`predict()` (*evalml.objectives.PrecisionWeighted method*), 103
`predict()` (*evalml.objectives.R2 method*), 123
`predict()` (*evalml.objectives.Recall method*), 105
`predict()` (*evalml.objectives.RecallMacro method*), 108
`predict()` (*evalml.objectives.RecallMicro method*), 106
`predict()` (*evalml.objectives.RecallWeighted method*), 109
`predict()` (*evalml.objectives.ROC method*), 120
`predict()` (*evalml.pipelines.components.LinearRegressor method*), 71
`predict()` (*evalml.pipelines.components.LogisticRegressionClassifier method*), 67
`predict()` (*evalml.pipelines.components.RandomForestClassifier method*), 69
`predict()` (*evalml.pipelines.components.RandomForestRegressor method*), 73
`predict()` (*evalml.pipelines.components.XGBoostClassifier method*), 70
`predict()` (*evalml.pipelines.LinearRegressionPipeline method*), 86
`predict()` (*evalml.pipelines.LogisticRegressionPipeline method*), 82
`predict()` (*evalml.pipelines.PipelineBase method*), 75
`predict()` (*evalml.pipelines.RFClassificationPipeline method*), 77
`predict()` (*evalml.pipelines.RFRegressionPipeline method*), 84
`predict()` (*evalml.pipelines.XGBoostPipeline method*), 79
`predict_proba()` (*evalml.pipelines.components.LinearRegressor method*), 71
`predict_proba()` (*evalml.pipelines.components.LogisticRegressionClassifier method*), 67
`predict_proba()` (*evalml.pipelines.components.RandomForestClassifier method*), 69
`predict_proba()` (*evalml.pipelines.components.RandomForestRegressor method*), 73
`predict_proba()` (*evalml.pipelines.components.XGBoostClassifier method*), 70
`predict_proba()` (*evalml.pipelines.LinearRegressionPipeline method*), 86
`predict_proba()` (*evalml.pipelines.LogisticRegressionPipeline method*), 82
`predict_proba()` (*evalml.pipelines.PipelineBase method*), 75
`predict_proba()` (*evalml.pipelines.RFClassificationPipeline method*), 78
`predict_proba()` (*evalml.pipelines.RFRegressionPipeline method*), 84
`predict_proba()` (*evalml.pipelines.XGBoostPipeline method*), 80
 ProblemTypes (*class in evalml.problem_types*), 133
`propose()` (*evalml.tuners.SKOptTuner method*), 134

R

R2 (class in *evalml.objectives*), 122

RandomForestClassifier (class in *evalml.pipelines.components*), 68

RandomForestRegressor (class in *evalml.pipelines.components*), 72

Recall (class in *evalml.objectives*), 104

RecallMacro (class in *evalml.objectives*), 107

RecallMicro (class in *evalml.objectives*), 105

RecallWeighted (class in *evalml.objectives*), 108

RFClassificationPipeline (class in *evalml.pipelines*), 76

RFClassifierSelectFromModel (class in *evalml.pipelines.components*), 62

RFRegressionPipeline (class in *evalml.pipelines*), 82

RFRegressorSelectFromModel (class in *evalml.pipelines.components*), 60

ROC (class in *evalml.objectives*), 119

S

save_pipeline() (in module *evalml.pipelines*), 73

score() (*evalml.objectives.AUC* method), 111

score() (*evalml.objectives.AUCMacro* method), 114

score() (*evalml.objectives.AUCMicro* method), 112

score() (*evalml.objectives.AUCWeighted* method), 115

score() (*evalml.objectives.ConfusionMatrix* method), 121

score() (*evalml.objectives.ExpVariance* method), 132

score() (*evalml.objectives.F1* method), 93

score() (*evalml.objectives.F1Macro* method), 96

score() (*evalml.objectives.F1Micro* method), 94

score() (*evalml.objectives.F1Weighted* method), 97

score() (*evalml.objectives.FraudCost* method), 89

score() (*evalml.objectives.LeadScoring* method), 91

score() (*evalml.objectives.LogLoss* method), 117

score() (*evalml.objectives.MAE* method), 125

score() (*evalml.objectives.MaxError* method), 131

score() (*evalml.objectives.MCC* method), 118

score() (*evalml.objectives.MedianAE* method), 129

score() (*evalml.objectives.MSE* method), 126

score() (*evalml.objectives.MSLE* method), 128

score() (*evalml.objectives.Precision* method), 99

score() (*evalml.objectives.PrecisionMacro* method), 102

score() (*evalml.objectives.PrecisionMicro* method), 100

score() (*evalml.objectives.PrecisionWeighted* method), 103

score() (*evalml.objectives.R2* method), 123

score() (*evalml.objectives.Recall* method), 105

score() (*evalml.objectives.RecallMacro* method), 108

score() (*evalml.objectives.RecallMicro* method), 106

score() (*evalml.objectives.RecallWeighted* method), 109

score() (*evalml.objectives.ROC* method), 120

score() (*evalml.pipelines.LinearRegressionPipeline* method), 86

score() (*evalml.pipelines.LogisticRegressionPipeline* method), 82

score() (*evalml.pipelines.PipelineBase* method), 76

score() (*evalml.pipelines.RFClassificationPipeline* method), 78

score() (*evalml.pipelines.RFRegressionPipeline* method), 84

score() (*evalml.pipelines.XGBoostPipeline* method), 80

SimpleImputer (class in *evalml.pipelines.components*), 63

SKOptTuner (class in *evalml.tuners*), 133

split_data() (in module *evalml.preprocessing*), 53

StandardScaler (class in *evalml.pipelines.components*), 65

supports_problem_type() (*evalml.objectives.AUC* class method), 111

supports_problem_type() (*evalml.objectives.AUCMacro* class method), 114

supports_problem_type() (*evalml.objectives.AUCMicro* class method), 113

supports_problem_type() (*evalml.objectives.AUCWeighted* class method), 116

supports_problem_type() (*evalml.objectives.ConfusionMatrix* class method), 122

supports_problem_type() (*evalml.objectives.ExpVariance* class method), 132

supports_problem_type() (*evalml.objectives.F1* class method), 93

supports_problem_type() (*evalml.objectives.F1Macro* class method), 96

supports_problem_type() (*evalml.objectives.F1Micro* class method), 95

supports_problem_type() (*evalml.objectives.F1Weighted* class method), 98

supports_problem_type() (*evalml.objectives.FraudCost* class method), 89

supports_problem_type() (*evalml.objectives.LeadScoring* class method), 91

supports_problem_type()

[\(evalml.objectives.LogLoss class method\), 117](#)
[supports_problem_type\(\) \(evalml.objectives.MAE class method\), 125](#)
[supports_problem_type\(\) \(evalml.objectives.MaxError class method\), 131](#)
[supports_problem_type\(\) \(evalml.objectives.MCC class method\), 119](#)
[supports_problem_type\(\) \(evalml.objectives.MedianAE class method\), 129](#)
[supports_problem_type\(\) \(evalml.objectives.MSE class method\), 126](#)
[supports_problem_type\(\) \(evalml.objectives.MSLE class method\), 128](#)
[supports_problem_type\(\) \(evalml.objectives.Precision class method\), 99](#)
[supports_problem_type\(\) \(evalml.objectives.PrecisionMacro class method\), 102](#)
[supports_problem_type\(\) \(evalml.objectives.PrecisionMicro class method\), 101](#)
[supports_problem_type\(\) \(evalml.objectives.PrecisionWeighted class method\), 104](#)
[supports_problem_type\(\) \(evalml.objectives.R2 class method\), 123](#)
[supports_problem_type\(\) \(evalml.objectives.Recall class method\), 105](#)
[supports_problem_type\(\) \(evalml.objectives.RecallMacro class method\), 108](#)
[supports_problem_type\(\) \(evalml.objectives.RecallMicro class method\), 107](#)
[supports_problem_type\(\) \(evalml.objectives.RecallWeighted class method\), 110](#)
[supports_problem_type\(\) \(evalml.objectives.ROC class method\), 120](#)

T

[transform\(\) \(evalml.pipelines.components.OneHotEncoder method\), 59](#)
[transform\(\) \(evalml.pipelines.components.RFClassifierSelectFromModel method\), 63](#)
[transform\(\) \(evalml.pipelines.components.RFRegressorSelectFromModel method\), 61](#)
[transform\(\) \(evalml.pipelines.components.SimpleImputer method\), 65](#)

X

[XGBoostClassifier \(class in evalml.pipelines.components\), 69](#)
[XGBoostPipeline \(class in evalml.pipelines\), 78](#)