

Appendix D: Default and Final Parameters of Estimators

Table 1 Gradient Boosting for Regression

	Default Parameters	Hyperparameters	Best Parameters
alpha	0.900000		0.900000
ccp_alpha	0.000000		0.000000
criterion	friedman_mse		friedman_mse
init			None
learning_rate	0.100000	[0.01 0.2075 0.405 0.6025 0.8]	0.207500
loss	squared_error	['squared_error', 'huber', 'quantile']	squared_error
max_depth	3	[2, 3, 4, 5, 6, None]	2
max_features			None
max_leaf_nodes			None
min_impurity_decrease	0.000000		0.000000
min_samples_leaf	1	[1, 2, 3, 4]	3
min_samples_split	2		2
min_weight_fraction_leaf	0.000000		0.000000
n_estimators	100	[50, 80, 100, 120, 150]	50
n_iter_no_change			None
random_state	42		42
subsample	1.000000		1.000000
tol	0.000100		0.000100
validation_fraction	0.100000		0.100000
verbose	0		0
warm_start	False		False

Default, hyper- and best parameters.

Table 2 K-nearest Neighbors Regression

	Default Parameters	Hyperparameters	Best Parameters
algorithm	auto	['ball_tree', 'kd_tree', 'brute', 'auto']	ball_tree
leaf_size	30	[5 10 15 20 25 30 35 40 45 50]	15
metric	minkowski	['minkowski', 'l1', 'l2']	minkowski
metric_params			None
n_jobs			None
n_neighbors	5	[5 13 21 29 37 45]	21
p	2	[1, 2]	1
weights	uniform	['uniform', 'distance']	uniform

Default, hyper- and best parameters.

Table 3 Random Forest Regressor

	Default Parameters	Hyperparameters	Best Parameters
bootstrap	True	[True]	True
ccp_alpha	0.000000		0.000000
criterion	squared_error	['squared_error', 'friedman_mse']	friedman_mse
max_depth		[3, 9, 16, 23, 30]	9
max_features	1.000000	['sqrt', 'log2', 0.6, 1.0]	log2
max_leaf_nodes			None
max_samples			None
min_impurity_decrease	0.000000		0.000000
min_samples_leaf	1	[1 4 7]	1
min_samples_split	2	[2, 4, 6]	4
min_weight_fraction_leaf	0.000000		0.000000
n_estimators	100	[50, 87, 125, 162, 200]	200
n_jobs			None
oob_score	False		False
random_state	42		42
verbose	0		0
warm_start	False		False

Default, hyper- and best parameters.

Table 4 Epsilon-Support Vector Regression

	Default Parameters	Hyperparameters	Best Parameters
C	1.000000	[0.01, 0.1, 1.0, 10, 100]	0.100000
cache_size	200	[10000]	10000
coef0	0.000000	[0.0, 0.1, 0.2]	0.100000
degree	3		3
epsilon	0.100000	[0.001, 0.01, 0.1, 1.0, 10]	0.100000
gamma	scale	['auto', 'scale']	scale
kernel	rbf	['rbf', 'sigmoid']	rbf
max_iter	-1		-1
shrinking	True	[True, False]	False
tol	0.001000	[0.0001, 0.001, 0.01, 0.1, 1.0]	0.010000
verbose	False		False

Default, hyper- and best parameters.

Table 5 Linear model fitted by minimizing a regularized empirical loss with stochastic gradient descent

	Default Parameters	Hyperparameters1	Hyperparameters2	Hyperparameters3	Hyperparameters4	Best Parameters
alpha	0.000100	[1e-05, 0.0001, 0.001, 0.01, 0.1]	[1e-05, 0.0001, 0.001, 0.01, 0.1]	[1e-05, 0.0001, 0.001, 0.01, 0.1]	[1e-05, 0.0001, 0.001, 0.01, 0.1]	0.000010
average	False					False
early_stopping	False					False
epsilon	0.100000					0.500000
eta0	0.010000	[0.01, 0.1, 0.5]	[0.01, 0.1, 0.5]	[0.01, 0.1, 0.5]	[0.01, 0.1, 0.5]	0.100000
fit_intercept	True	[0.001, 0.01, 0.1]	[0.001, 0.01, 0.1]	[0.001, 0.01, 0.1]	[0.001, 0.01, 0.1]	True
l1_ratio	0.150000			[0.1, 0.15, 0.2, 0.5]	[0.1, 0.15, 0.2, 0.5]	0.500000
learning_rate	invscaling	['constant', 'optimal', 'invscaling', 'adaptive']	['constant', 'optimal', 'invscaling', 'adaptive']	['constant', 'optimal', 'invscaling', 'adaptive']	['constant', 'optimal', 'invscaling', 'adaptive']	adaptive
loss	squared_error	['squared_error']	['squared_error']	['squared_error']	['squared_error']	squared_error
max_iter	100000000			['epsilon.insensitive', 'squared_epsilon.insensitive']	['epsilon.insensitive', 'squared_epsilon.insensitive']	100000000
n_iter_no_change	5					5
penalty	l2			['elasticnet']	['elasticnet']	elasticnet
power_t	0.250000	[None, 'l2', 'l1']	[None, 'l2', 'l1']	[0.1, 0.25, 0.5, 0.7]	[0.1, 0.25, 0.5, 0.7]	0.100000
random_state	42					42
shuffle	True					True
tol	0.001000	[0.0001, 0.001, 0.01, 0.1]	[0.0001, 0.001, 0.01, 0.1]	[0.0001, 0.001, 0.01, 0.1]	[0.0001, 0.001, 0.01, 0.1]	0.001000
validation_fraction	0.100000					0.100000
verbose	0					0
warm_start	False					False

Default, hyper- and best parameters.

Table 6 Multivariable ordinary least squares Linear Regression

	Default Parameters	Hyperparameters	Best Parameters
copy_X	True	[False, True]	False
fit_intercept	True	[False, True]	True
n_jobs			None
positive	False	[False, True]	True

Default, hyper- and best parameters.

Table 7 Multi-layer Perceptron Regressor

	Default Parameters	Hyperparameters1	Hyperparameters2	Hyperparameters3	Hyperparameters4	Best Parameters
activation	relu	['relu', 'tanh', 'logistic']	['relu', 'tanh', 'logistic']	['relu', 'tanh', 'logistic']	['relu', 'tanh', 'logistic']	tanh
alpha	0.000100	[0.0001, 0.001]	[0.0001, 0.001]	[0.0001, 0.001]	[0.0001, 0.001]	0.000100
batch_size	auto	['auto', 64, 128]	['auto', 64, 128]	['auto', 64, 128]	['auto', 64, 128]	128
beta.1	0.900000					0.900000
beta.2	0.999000					0.999000
early_stopping	False					True
epsilon	0.000000					0.000000
hidden_layer_sizes	(100,)	[(100,), (256, 128), (64,)]	[(100,), (256, 128), (64,)]	[(100,), (256, 128), (64,)]	[(100,), (256, 128), (64,)]	(256, 128)
learning_rate	constant	['constant', 'invscaling', 'adaptive']	['constant', 'invscaling', 'adaptive']	['constant', 'invscaling', 'adaptive']	['constant', 'invscaling', 'adaptive']	invscaling
learning_rate_init	0.001000	[0.001, 0.00001]	[0.001, 0.00001]	[0.001, 0.00001]	[0.001, 0.00001]	0.001000
max_fun	15000					15000
max_iter	200	[350]	[500]	[200, 350, 500]	[200, 350, 500]	350
momentum	0.900000	[0.5, 0.7, 0.9]	[0.5, 0.7, 0.9]			0.900000
n_iter_no_change	10					10
nesterovs_momentum	True					True
power_t	0.500000	[0.3, 0.5, 0.7]	[0.3, 0.5, 0.7]			0.700000
random_state	42					42
shuffle	True					True
solver	adam					adam
tol	0.000100					0.000100
validation_fraction	0.100000					0.100000
verbose	False					False
warm_start	False					False

Default, hyper- and best parameters.