

In [1]:

```
%load_ext watermark  
%watermark -a 'Sebastian Raschka' -u -d -v -p numpy,pandas,matplotlib,scipy,scikit-learn
```

UsageError: unrecognized arguments: Raschka'

In [4]:

```
from IPython.display import Image
```

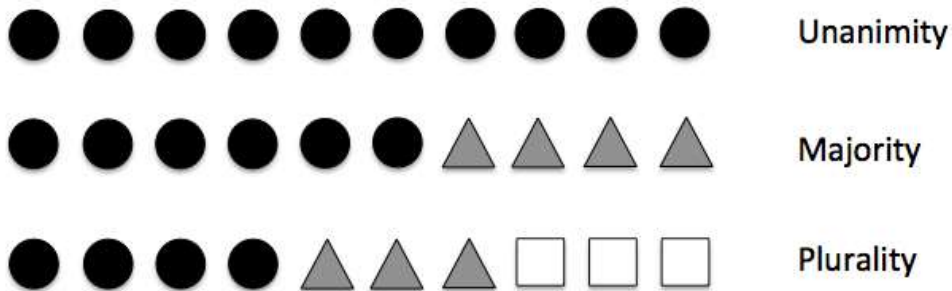
In [5]:

```
%matplotlib inline
```

In [6]:

```
Image(filename='07_01.png', width=500)
```

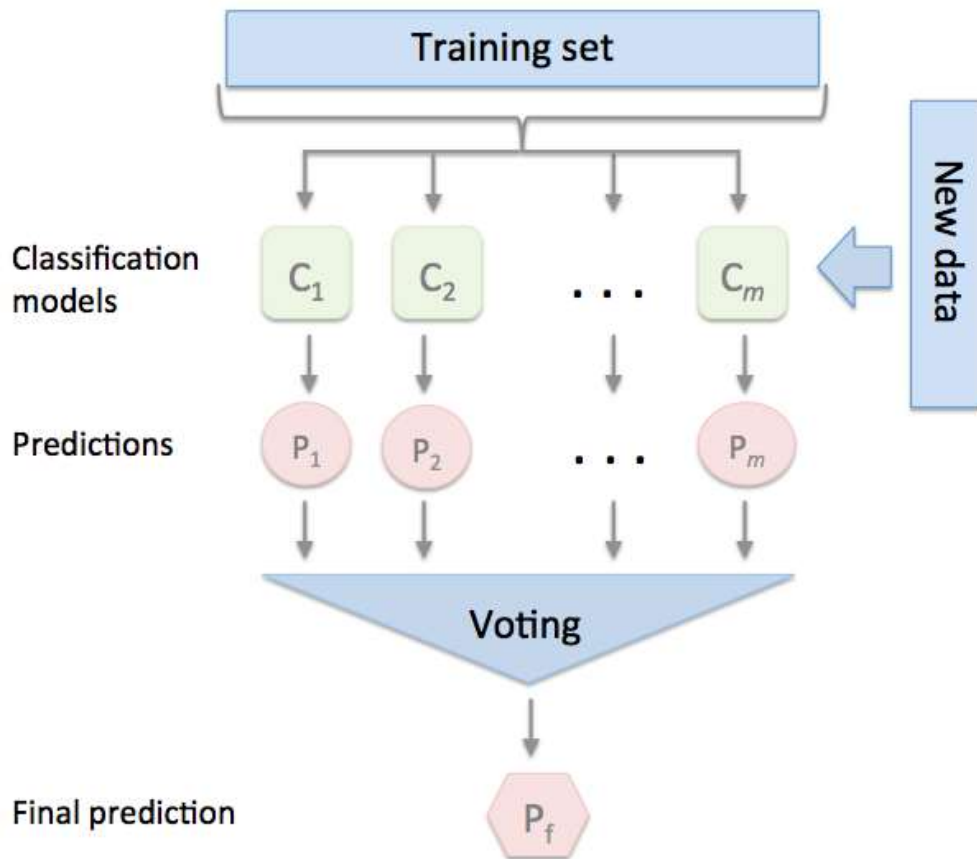
Out[6]:



In [7]:

```
Image(filename='07_02.png', width=500)
```

Out[7]:



In [8]:

```

from scipy.misc import comb
import math

def ensemble_error(n_classifier, error):
    k_start = math.ceil(n_classifier / 2.0)
    probs = [comb(n_classifier, k) * error**k * (1-error)**(n_classifier - k)
              for k in range(k_start, n_classifier + 1)]
    return sum(probs)

```

In [9]:

```

from scipy.misc import comb
import math

def ensemble_error(n_classifier, error):
    k_start = int(math.ceil(n_classifier / 2.0))
    probs = [comb(n_classifier, k) * error**k * (1-error)**(n_classifier - k)
              for k in range(k_start, n_classifier + 1)]
    return sum(probs)

```

In [10]:

```
ensemble_error(n_classifier=11, error=0.25)
```

Out[10]:

0.034327507019042969

In [11]:

```
import numpy as np

error_range = np.arange(0.0, 1.01, 0.01)
ens_errors = [ensemble_error(n_classifier=11, error=error)
               for error in error_range]
```

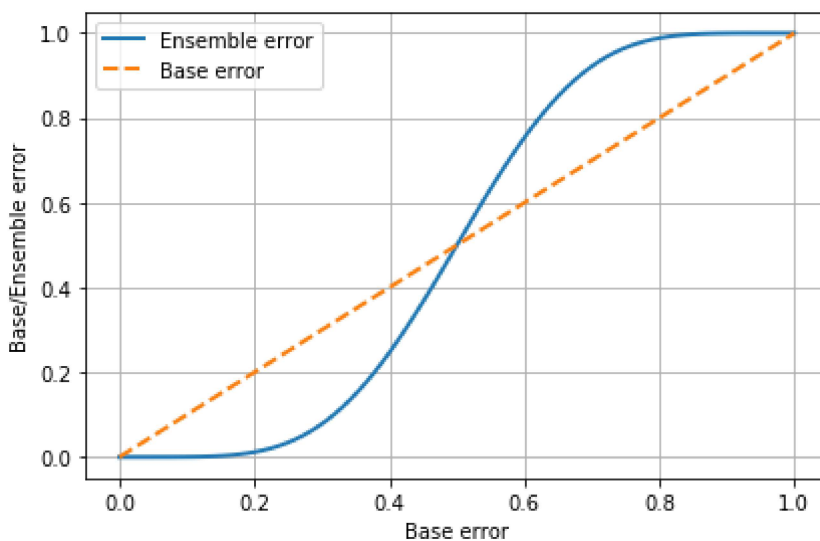
In [12]:

```
import matplotlib.pyplot as plt

plt.plot(error_range,
         ens_errors,
         label='Ensemble error',
         linewidth=2)

plt.plot(error_range,
         error_range,
         linestyle='--',
         label='Base error',
         linewidth=2)

plt.xlabel('Base error')
plt.ylabel('Base/Ensemble error')
plt.legend(loc='upper left')
plt.grid()
plt.tight_layout()
# plt.savefig('./figures/ensemble_err.png', dpi=300)
plt.show()
```



In [13]:

```
import numpy as np

np.argmax(np.bincount([0, 0, 1],
                      weights=[0.2, 0.2, 0.6]))
```

Out[13]:

1

In [14]:

```
ex = np.array([[0.9, 0.1],
               [0.8, 0.2],
               [0.4, 0.6]])

p = np.average(ex,
               axis=0,
               weights=[0.2, 0.2, 0.6])

p
```

Out[14]:

array([0.58, 0.42])

In [15]:

```
np.argmax(p)
```

Out[15]:

0

In [17]:

```
from sklearn.base import BaseEstimator
from sklearn.base import ClassifierMixin
from sklearn.preprocessing import LabelEncoder
from sklearn.externals import six
from sklearn.base import clone
from sklearn.pipeline import _name_estimators
import numpy as np
import operator

class MajorityVoteClassifier(BaseEstimator,
                           ClassifierMixin):
    """ A majority vote ensemble classifier

    Parameters
    -----
    classifiers : array-like, shape = [n_classifiers]
        Different classifiers for the ensemble

    vote : str, {'classlabel', 'probability'} (default='label')
        If 'classlabel' the prediction is based on the argmax of
        class labels. Else if 'probability', the argmax of
        the sum of probabilities is used to predict the class label
        (recommended for calibrated classifiers).

    weights : array-like, shape = [n_classifiers], optional (default=None)
```

If a list of `int` or `float` values are provided, the classifiers are weighted by importance; Uses uniform weights if `weights=None`.

"""

```
def __init__(self, classifiers, vote='classlabel', weights=None):

    self.classifiers = classifiers
    self.named_classifiers = {key: value for key, value
                              in _name_estimators(classifiers)}

    self.vote = vote
    self.weights = weights

def fit(self, X, y):
    """ Fit classifiers.

    Parameters
    -----
    X : {array-like, sparse matrix}, shape = [n_samples, n_features]
    Matrix of training samples.

    y : array-like, shape = [n_samples]
    Vector of target class labels.

    Returns
    -----
    self : object

    """

    if self.vote not in ('probability', 'classlabel'):
        raise ValueError("vote must be 'probability' or 'classlabel'"
                        "; got (vote=%r)"
                        % self.vote)

    if self.weights and len(self.weights) != len(self.classifiers):
        raise ValueError('Number of classifiers and weights must be equal'
                        '; got %d weights, %d classifiers'
                        % (len(self.weights), len(self.classifiers)))

    # Use LabelEncoder to ensure class labels start with 0, which
    # is important for np.argmax call in self.predict
    self.lablenc_ = LabelEncoder()
    self.lablenc_.fit(y)
    self.classes_ = self.lablenc_.classes_
    self.classifiers_ = []
    for clf in self.classifiers:
        fitted_clf = clone(clf).fit(X, self.lablenc_.transform(y))
        self.classifiers_.append(fitted_clf)
    return self

def predict(self, X):
    """ Predict class labels for X.

    Parameters
    -----
    X : {array-like, sparse matrix}, shape = [n_samples, n_features]
    Matrix of training samples.

    Returns
    -----
    maj_vote : array-like, shape = [n_samples]
    Predicted class labels.
```

```

"""
if self.vote == 'probability':
    maj_vote = np.argmax(self.predict_proba(X), axis=1)
else: # 'class/label' vote

    # Collect results from clf.predict calls
    predictions = np.asarray([clf.predict(X)
                              for clf in self.classifiers_]).T

    maj_vote = np.apply_along_axis(
        lambda x:
            np.argmax(np.bincount(x,
                                   weights=self.weights)),
        axis=1,
        arr=predictions)
maj_vote = self.lablenc_.inverse_transform(maj_vote)
return maj_vote

def predict_proba(self, X):
    """ Predict class probabilities for X.

    Parameters
    -----
    X : {array-like, sparse matrix}, shape = [n_samples, n_features]
        Training vectors, where n_samples is the number of samples and
        n_features is the number of features.

    Returns
    -----
    avg_proba : array-like, shape = [n_samples, n_classes]
        Weighted average probability for each class per sample.

    """
    probas = np.asarray([clf.predict_proba(X)
                          for clf in self.classifiers_])
    avg_proba = np.average(probas, axis=0, weights=self.weights)
    return avg_proba

def get_params(self, deep=True):
    """ Get classifier parameter names for GridSearch """
    if not deep:
        return super(MajorityVoteClassifier, self).get_params(deep=False)
    else:
        out = self.named_classifiers.copy()
        for name, step in six.iteritems(self.named_classifiers):
            for key, value in six.iteritems(step.get_params(deep=True)):
                out['%s__%s' % (name, key)] = value
        return out

```

In [47]:

```
from sklearn import datasets
from sklearn.preprocessing import StandardScaler
from sklearn.preprocessing import LabelEncoder
from sklearn.model_selection import train_test_split

iris = datasets.load_iris()
X, y = iris.data[50:, [1, 2]], iris.target[50:]
le = LabelEncoder()
y = le.fit_transform(y)

X_train, X_test, y_train, y_test = \
    train_test_split(X, y,
                    test_size=0.5,
                    random_state=1)
```

In [48]:

```
import numpy as np
from sklearn.linear_model import LogisticRegression
from sklearn.tree import DecisionTreeClassifier
from sklearn.neighbors import KNeighborsClassifier
from sklearn.pipeline import Pipeline
from sklearn.model_selection import cross_val_score

clf1 = LogisticRegression(penalty='l2',
                        C=0.001,
                        random_state=0)

clf2 = DecisionTreeClassifier(max_depth=1,
                            criterion='entropy',
                            random_state=0)

clf3 = KNeighborsClassifier(n_neighbors=1,
                          p=2,
                          metric='minkowski')

pipe1 = Pipeline([['sc', StandardScaler()],
                  ['clf', clf1]])
pipe3 = Pipeline([['sc', StandardScaler()],
                  ['clf', clf3]])

clf_labels = ['Logistic Regression', 'Decision Tree', 'KNN']

print('10-fold cross validation:\n')
for clf, label in zip([pipe1, clf2, pipe3], clf_labels):
    scores = cross_val_score(estimator=clf,
                            X=X_train,
                            y=y_train,
                            cv=10,
                            scoring='roc_auc')
    print("ROC AUC: %0.2f (+/- %0.2f) [%s]"
          % (scores.mean(), scores.std(), label))
```

10-fold cross validation:

```
ROC AUC: 0.92 (+/- 0.20) [Logistic Regression]
ROC AUC: 0.92 (+/- 0.15) [Decision Tree]
ROC AUC: 0.93 (+/- 0.10) [KNN]
```


In [49]:

```
# Majority Rule (hard) Voting

mv_clf = MajorityVoteClassifier(classifiers=[pipe1, clf2, pipe3])

clf_labels += ['Majority Voting']
all_clf = [pipe1, clf2, pipe3, mv_clf]

for clf, label in zip(all_clf, clf_labels):
    scores = cross_val_score(estimator=clf,
                              X=X_train,
                              y=y_train,
                              cv=10,
                              scoring='roc_auc')
    print("ROC AUC: %0.2f (+/- %0.2f) [%s]"
          % (scores.mean(), scores.std(), label))
```

ROC AUC: 0.92 (+/- 0.20) [Logistic Regression]

ROC AUC: 0.92 (+/- 0.15) [Decision Tree]

ROC AUC: 0.93 (+/- 0.10) [KNN]

ROC AUC: 0.97 (+/- 0.10) [Majority Voting]

In [50]:

```
from sklearn.metrics import roc_curve
from sklearn.metrics import auc

colors = ['black', 'orange', 'blue', 'green']
linestyles = [':', '--', '-.', '-']
for clf, label, clr, ls in
    in zip(all_clf,
           clf_labels, colors, linestyles):

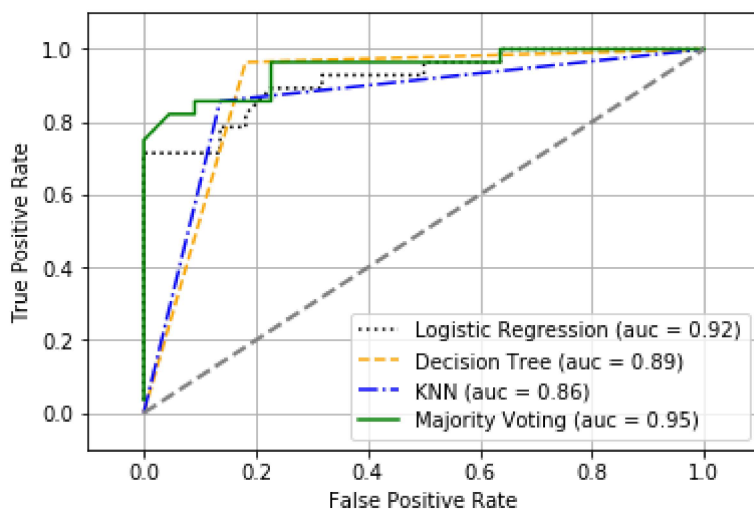
    # assuming the label of the positive class is 1
    y_pred = clf.fit(X_train,
                     y_train).predict_proba(X_test)[:, 1]
    fpr, tpr, thresholds = roc_curve(y_true=y_test,
                                     y_score=y_pred)

    roc_auc = auc(x=fpr, y=tpr)
    plt.plot(fpr, tpr,
             color=clr,
             linestyle=ls,
             label='%s (auc = %0.2f)' % (label, roc_auc))

plt.legend(loc='lower right')
plt.plot([0, 1], [0, 1],
         linestyle='--',
         color='gray',
         linewidth=2)

plt.xlim([-0.1, 1.1])
plt.ylim([-0.1, 1.1])
plt.grid()
plt.xlabel('False Positive Rate')
plt.ylabel('True Positive Rate')

# plt.tight_layout()
# plt.savefig('./figures/roc.png', dpi=300)
plt.show()
```



In [51]:

```
sc = StandardScaler()
X_train_std = sc.fit_transform(X_train)
```

In [52]:

```

from itertools import product

all_clf = [pipe1, clf2, pipe3, mv_clf]

x_min = X_train_std[:, 0].min() - 1
x_max = X_train_std[:, 0].max() + 1
y_min = X_train_std[:, 1].min() - 1
y_max = X_train_std[:, 1].max() + 1

xx, yy = np.meshgrid(np.arange(x_min, x_max, 0.1),
                     np.arange(y_min, y_max, 0.1))

f, axarr = plt.subplots(nrows=2, ncols=2,
                       sharex='col',
                       sharey='row',
                       figsize=(7, 5))

for idx, clf, tt in zip(product([0, 1], [0, 1]),
                       all_clf, clf_labels):
    clf.fit(X_train_std, y_train)

    Z = clf.predict(np.c_[xx.ravel(), yy.ravel()])
    Z = Z.reshape(xx.shape)

    axarr[idx[0], idx[1]].contourf(xx, yy, Z, alpha=0.3)

    axarr[idx[0], idx[1]].scatter(X_train_std[y_train==0, 0],
                                  X_train_std[y_train==0, 1],
                                  c='blue',
                                  marker='^',
                                  s=50)

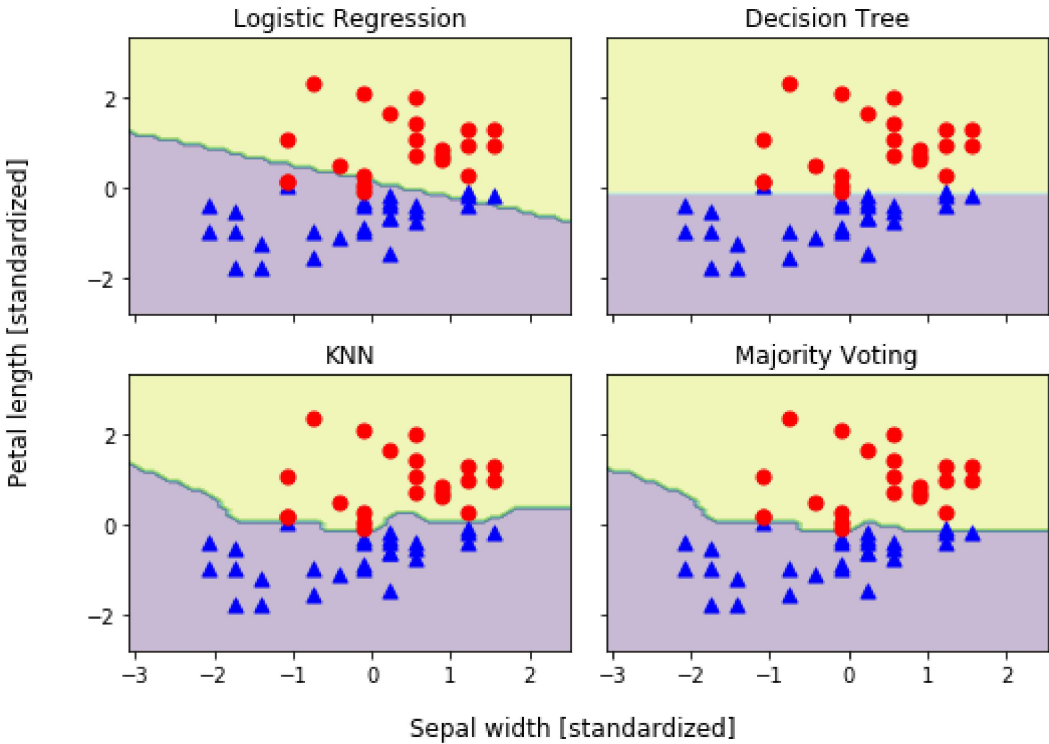
    axarr[idx[0], idx[1]].scatter(X_train_std[y_train==1, 0],
                                  X_train_std[y_train==1, 1],
                                  c='red',
                                  marker='o',
                                  s=50)

    axarr[idx[0], idx[1]].set_title(tt)

plt.text(-3.5, -4.5,
        s='Sepal width [standardized]',
        ha='center', va='center', fontsize=12)
plt.text(-10.5, 4.5,
        s='Petal length [standardized]',
        ha='center', va='center',
        fontsize=12, rotation=90)

plt.tight_layout()
# plt.savefig('./figures/voting_panel', bbox_inches='tight', dpi=300)
plt.show()

```



In [53]:

```
mv_clf.get_params()
```

Out[53]:

```
{'decisiontreeclassifier': DecisionTreeClassifier(class_weight=None, criterion='entropy', max_depth=1,
max_features=None, max_leaf_nodes=None,
min_impurity_split=1e-07, min_samples_leaf=1,
min_samples_split=2, min_weight_fraction_leaf=0.0,
presort=False, random_state=0, splitter='best'),
'decisiontreeclassifier__class_weight': None,
'decisiontreeclassifier__criterion': 'entropy',
'decisiontreeclassifier__max_depth': 1,
'decisiontreeclassifier__max_features': None,
'decisiontreeclassifier__max_leaf_nodes': None,
'decisiontreeclassifier__min_impurity_split': 1e-07,
'decisiontreeclassifier__min_samples_leaf': 1,
'decisiontreeclassifier__min_samples_split': 2,
'decisiontreeclassifier__min_weight_fraction_leaf': 0.0,
'decisiontreeclassifier__presort': False,
'decisiontreeclassifier__random_state': 0,
'decisiontreeclassifier__splitter': 'best',
'pipeline-1': Pipeline(steps=[['sc', StandardScaler(copy=True, with_mean=True, with_std=True)], ['clf', LogisticRegression(C=0.001, class_weight=None, dual=False, fit_intercept=True,
intercept_scaling=1, max_iter=100, multi_class='ovr', n_jobs=1,
penalty='l2', random_state=0, solver='liblinear', tol=0.0001,
verbose=0, warm_start=False)]],
'pipeline-1__clf': LogisticRegression(C=0.001, class_weight=None, dual=False, fit_intercept=True,
intercept_scaling=1, max_iter=100, multi_class='ovr', n_jobs=1,
penalty='l2', random_state=0, solver='liblinear', tol=0.0001,
verbose=0, warm_start=False),
'pipeline-1__clf__C': 0.001,
'pipeline-1__clf__class_weight': None,
'pipeline-1__clf__dual': False,
'pipeline-1__clf__fit_intercept': True,
'pipeline-1__clf__intercept_scaling': 1,
'pipeline-1__clf__max_iter': 100,
'pipeline-1__clf__multi_class': 'ovr',
'pipeline-1__clf__n_jobs': 1,
'pipeline-1__clf__penalty': 'l2',
'pipeline-1__clf__random_state': 0,
'pipeline-1__clf__solver': 'liblinear',
'pipeline-1__clf__tol': 0.0001,
'pipeline-1__clf__verbose': 0,
'pipeline-1__clf__warm_start': False,
'pipeline-1__sc': StandardScaler(copy=True, with_mean=True, with_std=True),
'pipeline-1__sc__copy': True,
'pipeline-1__sc__with_mean': True,
'pipeline-1__sc__with_std': True,
'pipeline-1__steps': [['sc',
StandardScaler(copy=True, with_mean=True, with_std=True)],
['clf',
LogisticRegression(C=0.001, class_weight=None, dual=False, fit_intercept=True,
intercept_scaling=1, max_iter=100, multi_class='ovr', n_jobs=1,
penalty='l2', random_state=0, solver='liblinear', tol=0.0001,
verbose=0, warm_start=False)]],
'pipeline-2': Pipeline(steps=[['sc', StandardScaler(copy=True, with_mean=True, with_std=True)], ['clf', KNeighborsClassifier(algorithm='auto', leaf_size=30, metric='minkowski',
metric_params=None, n_jobs=1, n_neighbors=1, p=2,
weights='uniform')]]),
'pipeline-2__clf': KNeighborsClassifier(algorithm='auto', leaf_size=30, metric='minkowski',
metric_params=None, n_jobs=1, n_neighbors=1, p=2,
weights='uniform')])
```

```

pipeline-2__clf__ = KNeighborsClassifier(algorithm='auto', leaf_size=30, metric='minkowski',
                                         metric_params=None, n_jobs=1, n_neighbors=1, p=2,
                                         weights='uniform'),
'pipeline-2__clf__algorithm': 'auto',
'pipeline-2__clf__leaf_size': 30,
'pipeline-2__clf__metric': 'minkowski',
'pipeline-2__clf__metric_params': None,
'pipeline-2__clf__n_jobs': 1,
'pipeline-2__clf__n_neighbors': 1,
'pipeline-2__clf__p': 2,
'pipeline-2__clf__weights': 'uniform',
'pipeline-2__sc__': StandardScaler(copy=True, with_mean=True, with_std=True),
'pipeline-2__sc__copy': True,
'pipeline-2__sc__with_mean': True,
'pipeline-2__sc__with_std': True,
'pipeline-2__steps': [['sc',
                       StandardScaler(copy=True, with_mean=True, with_std=True)],
                      ['clf',
                       KNeighborsClassifier(algorithm='auto', leaf_size=30, metric='minkowski',
                                           metric_params=None, n_jobs=1, n_neighbors=1, p=2,
                                           weights='uniform')]]]

```

In [54]:

```

from sklearn.grid_search import GridSearchCV

params = {'decisiontreeclassifier__max_depth': [1, 2],
          'pipeline-1__clf__C': [0.001, 0.1, 100.0]}

grid = GridSearchCV(estimator=mv_clf,
                    param_grid=params,
                    cv=10,
                    scoring='roc_auc')
grid.fit(X_train, y_train)

for params, mean_score, scores in grid.grid_scores_:
    print("%0.3f+/-%0.2f %r"
          % (mean_score, scores.std() / 2.0, params))

```

```

0.967+/-0.05 {'decisiontreeclassifier__max_depth': 1, 'pipeline-1__clf__C': 0.001}
0.967+/-0.05 {'decisiontreeclassifier__max_depth': 1, 'pipeline-1__clf__C': 0.1}
1.000+/-0.00 {'decisiontreeclassifier__max_depth': 1, 'pipeline-1__clf__C': 100.0}
0.967+/-0.05 {'decisiontreeclassifier__max_depth': 2, 'pipeline-1__clf__C': 0.001}
0.967+/-0.05 {'decisiontreeclassifier__max_depth': 2, 'pipeline-1__clf__C': 0.1}
1.000+/-0.00 {'decisiontreeclassifier__max_depth': 2, 'pipeline-1__clf__C': 100.0}

```

In [55]:

```

print('Best parameters: %s' % grid.best_params_)
print('Accuracy: %.2f' % grid.best_score_)

```

```

Best parameters: {'decisiontreeclassifier__max_depth': 1, 'pipeline-1__clf__C': 10
0.0}
Accuracy: 1.00

```

In [56]:

```
grid.best_estimator_.classifiers
```

Out[56]:

```
[Pipeline(steps=[('sc', StandardScaler(copy=True, with_mean=True, with_std=True)),
('clf', LogisticRegression(C=100.0, class_weight=None, dual=False, fit_intercept=True,
    intercept_scaling=1, max_iter=100, multi_class='ovr', n_jobs=1,
    penalty='l2', random_state=0, solver='liblinear', tol=0.0001,
    verbose=0, warm_start=False))]),
DecisionTreeClassifier(class_weight=None, criterion='entropy', max_depth=1,
    max_features=None, max_leaf_nodes=None,
    min_impurity_split=1e-07, min_samples_leaf=1,
    min_samples_split=2, min_weight_fraction_leaf=0.0,
    presort=False, random_state=0, splitter='best'),
Pipeline(steps=[('sc', StandardScaler(copy=True, with_mean=True, with_std=True)),
('clf', KNeighborsClassifier(algorithm='auto', leaf_size=30, metric='minkowski',
    metric_params=None, n_jobs=1, n_neighbors=1, p=2,
    weights='uniform'))])]
```

In [75]:

```
mv_clf = grid.best_estimator_
```

In [76]:

```
mv_clf.set_params(**grid.best_estimator_.get_params())
```

Out[76]:

```
MajorityVoteClassifier(classifiers=[Pipeline(steps=[('sc', StandardScaler(copy=True,
    with_mean=True, with_std=True)), ('clf', LogisticRegression(C=100.0, class_weight=None,
    dual=False, fit_intercept=True,
    intercept_scaling=1, max_iter=100, multi_class='ovr', n_jobs=1,
    penalty='l2', random_state=0, solver='liblinear', tol=0.0001,
    verbose=0, warm_start=False)),
    Pipeline(steps=[('sc', StandardScaler(copy=True, with_mean=True, with_std=True)),
    ('clf', KNeighborsClassifier(algorithm='auto', leaf_size=30, metric='minkowski',
    metric_params=None, n_jobs=1, n_neighbors=1, p=2,
    weights='uniform'))])]),
    vote='classlabel', weights=None)
```

In [77]:

```
mv_clf
```

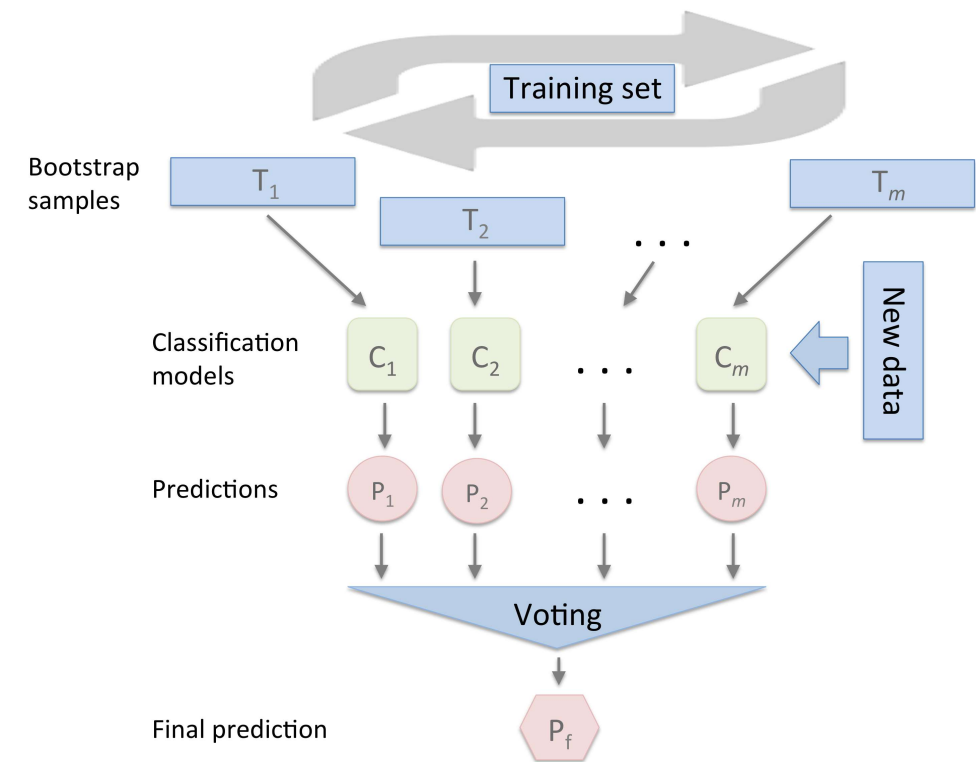
Out[77]:

```
MajorityVoteClassifier(classifiers=[Pipeline(steps=[('sc', StandardScaler(copy=True,
    with_mean=True, with_std=True)), ('clf', LogisticRegression(C=100.0, class_weight=None,
    dual=False, fit_intercept=True,
    intercept_scaling=1, max_iter=100, multi_class='ovr', n_jobs=1,
    penalty='l2', random_state=0, solver='liblinear', tol=0.0001,
    verbose=0, warm_start=False)),
    Pipeline(steps=[('sc', StandardScaler(copy=True, with_mean=True, with_std=True)),
    ('clf', KNeighborsClassifier(algorithm='auto', leaf_size=30, metric='minkowski',
    metric_params=None, n_jobs=1, n_neighbors=1, p=2,
    weights='uniform'))])]),
    vote='classlabel', weights=None)
```


In [78]:

```
Image(filename='07_06.png', width=500)
```

Out[78]:



In [63]:

```
Image(filename='07_07.png', width=400)
```

Out[63]:

Sample indices	Bagging round 1	Bagging round 2	...
1	2	7	...
2	2	3	...
3	1	2	...
4	3	1	...
5	7	1	...
6	2	7	...
7	4	7	...

Below the table, three brackets group the columns 'Bagging round 1', 'Bagging round 2', and '...'. Arrows point from these brackets to labels C_1 , C_2 , and C_m respectively, indicating that these columns represent the predictions of the individual models.

In [64]:

```
import pandas as pd

df_wine = pd.read_csv('https://archive.ics.uci.edu/ml/'
                     'machine-learning-databases/wine/wine.data',
                     header=None)

df_wine.columns = ['Class label', 'Alcohol', 'Malic acid', 'Ash',
                  'Alcalinity of ash', 'Magnesium', 'Total phenols',
                  'Flavanoids', 'Nonflavanoid phenols', 'Proanthocyanins',
                  'Color intensity', 'Hue', 'OD280/OD315 of diluted wines',
                  'Proline']

# drop 1 class
df_wine = df_wine[df_wine['Class label'] != 1]

y = df_wine['Class label'].values
X = df_wine[['Alcohol', 'Hue']].values
```

In [65]:

```
from sklearn.preprocessing import LabelEncoder
from sklearn.cross_validation import train_test_split

le = LabelEncoder()
y = le.fit_transform(y)

X_train, X_test, y_train, y_test = \
    train_test_split(X, y,
                    test_size=0.40,
                    random_state=1)
```

In [66]:

```
from sklearn.ensemble import BaggingClassifier
from sklearn.tree import DecisionTreeClassifier

tree = DecisionTreeClassifier(criterion='entropy',
                             max_depth=None,
                             random_state=1)

bag = BaggingClassifier(base_estimator=tree,
                       n_estimators=500,
                       max_samples=1.0,
                       max_features=1.0,
                       bootstrap=True,
                       bootstrap_features=False,
                       n_jobs=1,
                       random_state=1)
```

In [67]:

```
from sklearn.metrics import accuracy_score

tree = tree.fit(X_train, y_train)
y_train_pred = tree.predict(X_train)
y_test_pred = tree.predict(X_test)

tree_train = accuracy_score(y_train, y_train_pred)
tree_test = accuracy_score(y_test, y_test_pred)
print('Decision tree train/test accuracies %.3f/%.3f'
      % (tree_train, tree_test))

bag = bag.fit(X_train, y_train)
y_train_pred = bag.predict(X_train)
y_test_pred = bag.predict(X_test)

bag_train = accuracy_score(y_train, y_train_pred)
bag_test = accuracy_score(y_test, y_test_pred)
print('Bagging train/test accuracies %.3f/%.3f'
      % (bag_train, bag_test))
```

Decision tree train/test accuracies 1.000/0.833

Bagging train/test accuracies 1.000/0.896

In [68]:

```
import numpy as np
import matplotlib.pyplot as plt

x_min = X_train[:, 0].min() - 1
x_max = X_train[:, 0].max() + 1
y_min = X_train[:, 1].min() - 1
y_max = X_train[:, 1].max() + 1

xx, yy = np.meshgrid(np.arange(x_min, x_max, 0.1),
                     np.arange(y_min, y_max, 0.1))

f, axarr = plt.subplots(nrows=1, ncols=2,
                       sharex='col',
                       sharey='row',
                       figsize=(8, 3))

for idx, clf, tt in zip([0, 1],
                       [tree, bag],
                       ['Decision Tree', 'Bagging']):
    clf.fit(X_train, y_train)

    Z = clf.predict(np.c_[xx.ravel(), yy.ravel()])
    Z = Z.reshape(xx.shape)

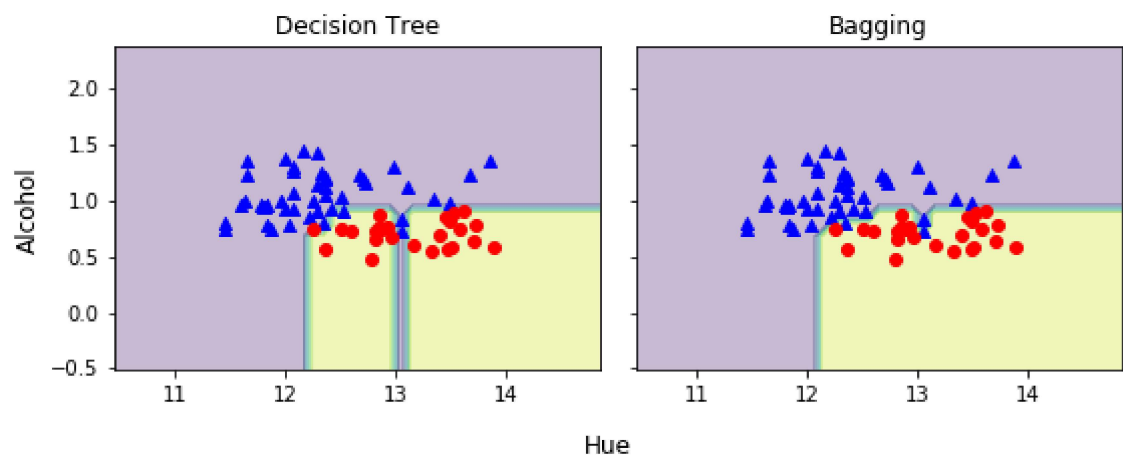
    axarr[idx].contourf(xx, yy, Z, alpha=0.3)
    axarr[idx].scatter(X_train[y_train == 0, 0],
                      X_train[y_train == 0, 1],
                      c='blue', marker='^')

    axarr[idx].scatter(X_train[y_train == 1, 0],
                      X_train[y_train == 1, 1],
                      c='red', marker='o')

    axarr[idx].set_title(tt)

axarr[0].set_ylabel('Alcohol', fontsize=12)
plt.text(10.2, -1.2,
        s='Hue',
        ha='center', va='center', fontsize=12)

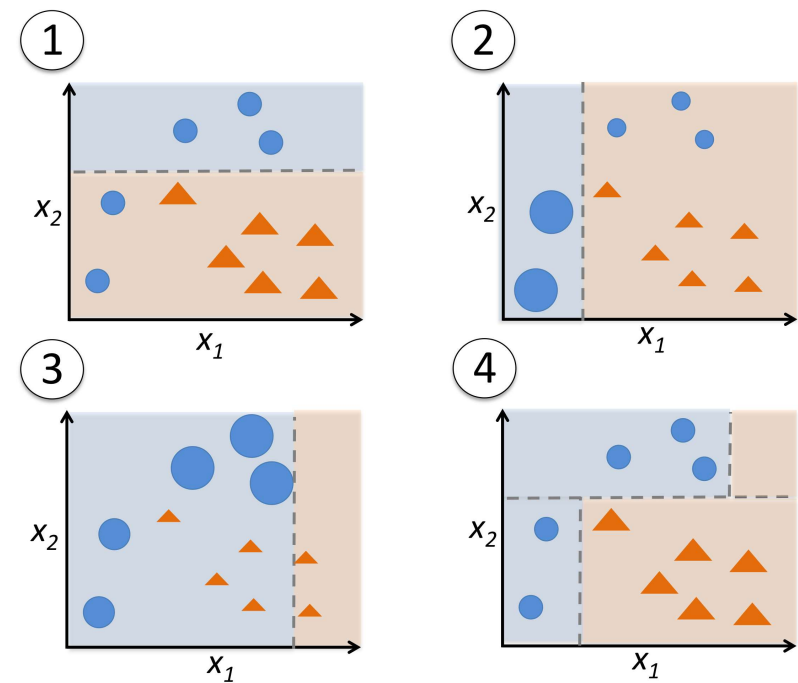
plt.tight_layout()
# plt.savefig('./figures/bagging_region.png',
#             dpi=300,
#             bbox_inches='tight')
plt.show()
```



In [69]:

```
Image(filename='07_09.png', width=400)
```

Out[69]:



In [71]:

```
Image(filename='07_10.png', width=500)
```

Out[71]:

Sample indices	x	y	Weights	$\hat{y}(x \leq 3.0)?$	Correct?	Updated weights
1	1.0	1	0.1	1	Yes	0.072
2	2.0	1	0.1	1	Yes	0.072
3	3.0	1	0.1	1	Yes	0.072
4	4.0	-1	0.1	-1	Yes	0.072
5	5.0	-1	0.1	-1	Yes	0.072
6	6.0	-1	0.1	-1	Yes	0.072
7	7.0	1	0.1	-1	No	0.167
8	8.0	1	0.1	-1	No	0.167
9	9.0	1	0.1	-1	No	0.167
10	10.0	-1	0.1	-1	Yes	0.072

In [72]:

```
from sklearn.ensemble import AdaBoostClassifier

tree = DecisionTreeClassifier(criterion='entropy',
                              max_depth=1,
                              random_state=0)

ada = AdaBoostClassifier(base_estimator=tree,
                          n_estimators=500,
                          learning_rate=0.1,
                          random_state=0)
```

In [73]:

```
tree = tree.fit(X_train, y_train)
y_train_pred = tree.predict(X_train)
y_test_pred = tree.predict(X_test)

tree_train = accuracy_score(y_train, y_train_pred)
tree_test = accuracy_score(y_test, y_test_pred)
print('Decision tree train/test accuracies %.3f/%.3f'
      % (tree_train, tree_test))

ada = ada.fit(X_train, y_train)
y_train_pred = ada.predict(X_train)
y_test_pred = ada.predict(X_test)

ada_train = accuracy_score(y_train, y_train_pred)
ada_test = accuracy_score(y_test, y_test_pred)
print('AdaBoost train/test accuracies %.3f/%.3f'
      % (ada_train, ada_test))
```

Decision tree train/test accuracies 0.845/0.854

AdaBoost train/test accuracies 1.000/0.875