



ROHINI COLLEGE OF ENGINEERING AND TECHNOLOGY
AUTONOMOUS INSTITUTION

Approved by AICTE & Affiliated to Anna University
NBA Accredited for BE (ECE, EEE, MECH) | Accredited by NAAC with A+ Grade
Anjugramam - Kanyakumari Main Road, Palkulam, Variyoor P.O. - 629 401, Kanyakumari District.

DEPARTMENT OF AGRICULTURAL ENGINEERING

AI3701 – REMOTE SENSING AND GEOGRAPHICAL INFORMATION SYATEM

UNIT 4 DATA INPUT AND ANALYSIS

4.2 IMAGE CLASSIFICATION

PREPARED BY

MRS. NITHYA.K

ASSISTANT PROFESSOR/AGRI



4.2 Image Classification

Image classification is a procedure to automatically categorize all pixels in an image of a terrain into land cover classes. Normally, multispectral data are used to perform the classification of the spectral pattern present within the data for each pixel is used as the numerical basis for categorization. This concept is dealt under the broad subject, namely, Pattern Recognition. Spectral pattern recognition refers to the family of classification procedures that utilises this pixel-by-pixel spectral information as the basis for automated land cover classification. Spatial pattern recognition involves the categorization of image pixels on the basis of the spatial relationship with pixels surrounding them. Image classification techniques are grouped into two types, namely supervised and unsupervised. The classification process may also include features, such as, land surface elevation and the soil type that are not derived from the image.

A pattern is thus a set of measurements on the chosen features for the individual to be classified. The classification process may therefore be considered a form of pattern recognition, that is, the identification of the pattern associated with each pixel position in an image in terms of the characteristics of the objects or on the earth's surface.

Supervised Classification

A supervised classification algorithm requires a training sample for each class, that is, a collection of data points known to have come from the class of interest. The classification is thus based on how "close" a point to be classified is to each training sample. We shall not attempt to define the word "close" other than to say that both geometric and statistical distance measures are used in practical pattern recognition algorithms. The training samples are representative of the known classes of interest to the analyst. Classification methods that relay on use of training patterns are called supervised classification methods.

The three basic steps involved in a typical supervised classification procedure are as follows :

(i) Training stage: The analyst identifies representative training areas and develops numerical descriptions of the spectral signatures of each land cover type of interest in the scene.

(ii) The classification stage: Each pixel in the image data set is categorized into the land cover class it most closely resembles. If the pixel is insufficiently similar to any training data set it is usually labeled 'Unknown'.

(iii) The output stage: The results may be used in a number of different ways. Three typical forms of output products are thematic maps, tables and digital data files which become input data for GIS. The output of image classification becomes input for GIS for spatial analysis of the terrain. Figure depicts the flow of operations to be performed during image classification of remotely sensed data of an area which ultimately leads to create database as an input for GIS. Plate 6 shows the land use/ land cover colour coded image, which is an output of image classification.

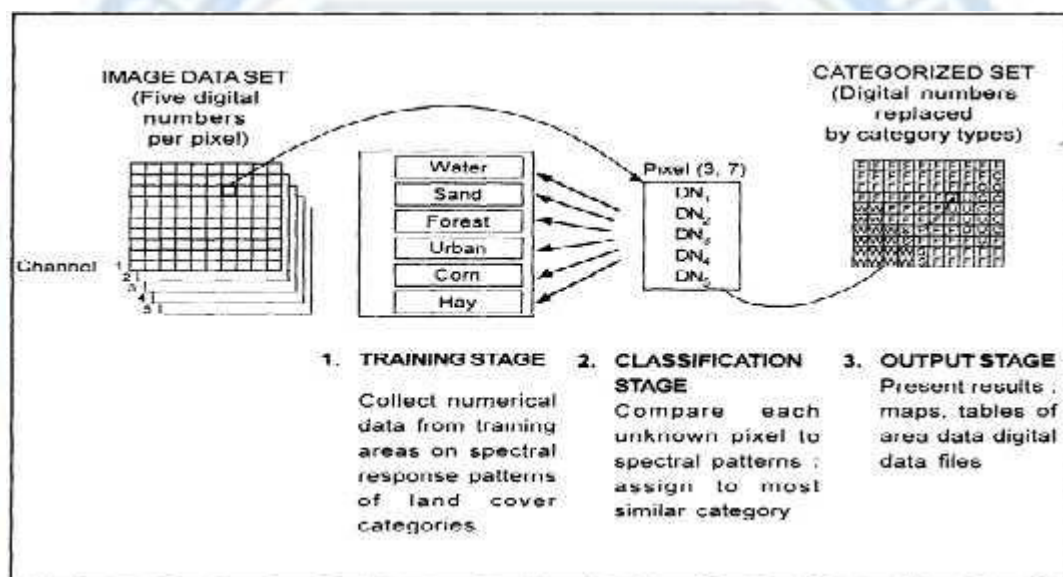


Figure Basic steps supervised classification

There are a number of powerful supervised classifiers based on the statistics, which are commonly, used for various applications. A few of them are a minimum distance to means method, average distance method, parallelepiped method, maximum likelihood method, modified maximum likelihood method, bayesian's method, decision tree classification, and discriminant functions. The principles and working algorithms of all these supervised classifiers are available in almost all standard books on remote sensing and so details are not provided here. Since all the supervised classification methods use training data samples, it is more appropriate to consider some of the fundamental characteristics of training data.

Training Dataset

A training dataset is a set of measurements (points from an image) whose category membership is known by the analyst. This set must be selected based on additional information derived from maps, field surveys, aerial photographs, and analyst's knowledge of usual spectral signatures of different cover classes. Selecting a good set of training points is one of the most critical aspects of the classification procedure. These guidelines are as following:

(i) Select sufficient number of points for each class. If each measurement vector has N features, then select $N+1$ points per class and the practical minimum is $10*N$ per class. If the class shows a lot of variability (the scatter plot showing considerable spreading or scatter among training points), select a larger number of points, subject to practical limits of time, effort and expense. The more the training points, the better the "extra points" to evaluate the accuracy of the classifier.

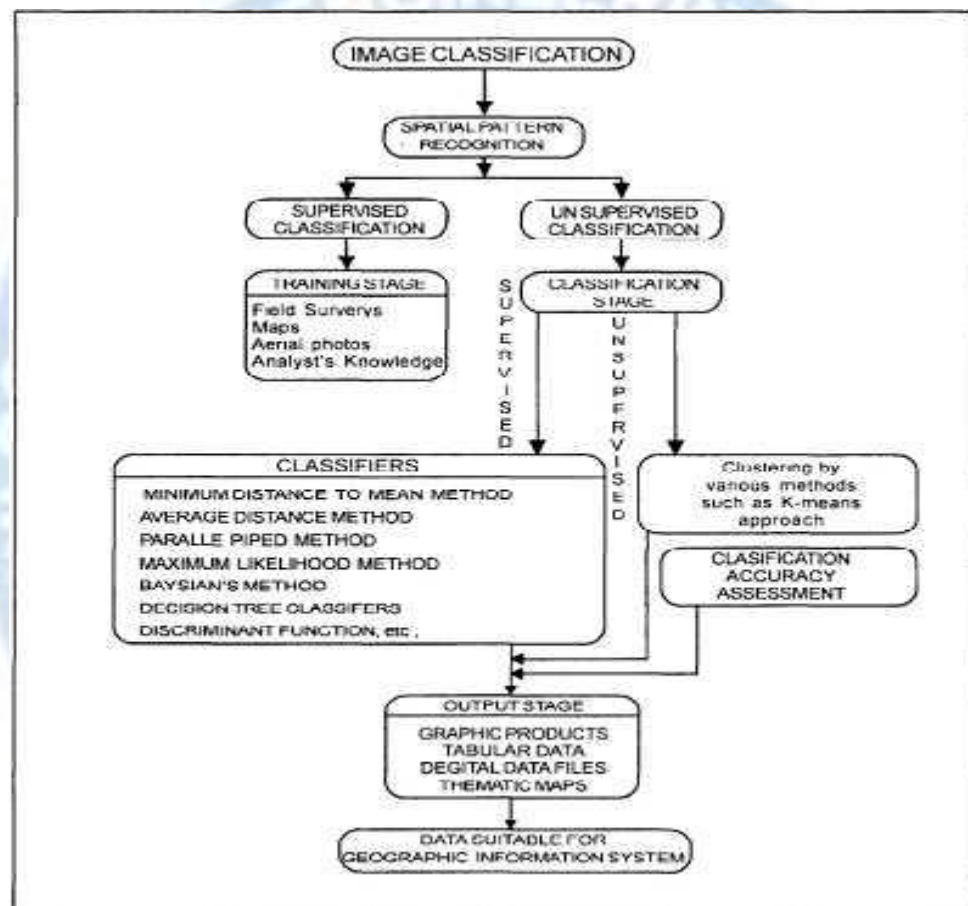


Figure Flow Chart showing Image Classification

The more the points, the more accurate the classification will be.

(ii) Select training data sets which are representative of the classes of interest that show both typical average feature values and a typical degree of variability. For each class, select several training areas on the image, instead of just one. Each training area should contain a moderately large number of pixels. Pick training areas from seemingly heterogeneous or appearing regions. Pick training areas that are widely and spatially dispersed, across the full image. For each class, select the training areas which are uniformly distributed across the image and with high density.

(iii) Check that selected areas have unimodal distributions (histograms). A bimodal histogram suggests that pixels from two different classes may be included in the training sample.

(iv) Select training sets (physically) using a computer-based classification system:

Poorest method: Using coordinates of training points or training regions directly.

Better method: using joystick, trackball, light pen, directly on the image.

For example. EASI/PACE : The program should show the histograms, mean and standard deviations for each region selected, and for each class in total.

The Program should allow to iterate; do classification using one set of training points, then come back and modify training sets and class definitions without starting all over again. There should be options to combine classes from previous classification.

(v) The program should allow one to designate half of the points as training points, and the other half to test the accuracy of the trained classifier. Before it is used, the training set should be evaluated by examining scatterplots and/or histograms for each class. It should show unimodal distributions, hopefully approximating normal distributions. If not unimodal, one may want to select new training sets. After the discriminant functions and the classification rule is derived, accuracy must be tested.

Two acceptable techniques which are commonly used are:

(a) Designate a randomly selected half of the training points as test points, before developing classifier. Use the other half for training. Then classify the half of the data not used for training. Develop contingency table (confusion matrix) to indicate probability of error in each class. This procedure is actually a measure of the consistency of the classifier.

(b) Randomly select a set of pixel regions from the image of an unknown class. Classify them using the discriminant function and rules developed from the training set. Then verify the correctness of the classification (again with a confusion matrix) by checking the identity of these regions using external information sources like maps and aerial photos.

(vi) Separability of classes: So far, we have looked at an ideal situation where there is no overlap between different classes. In reality the classes are likely to overlap. It can be seen that the less the overlap between classes the lower the chance of misclassifying a given pixel. Classes that have little overlap is said to be highly separable.

Unsupervised Classification

Unsupervised classification algorithms do not compare points to be classified with training data. Rather, unsupervised algorithms examine a large number of unknown data vectors and divide them into classes based on properties inherent to the data themselves. The classes that result stem from differences observed in the data. In particular, use is made of the notion that data vectors within a class should be in some sense mutually close together in the measurement space, whereas data vectors in different classes should be comparatively well separated. If the components of the data vectors represent the responses in different spectral bands, the resulting classes might be referred to as spectral classes, as opposed to information classes, which represent the ground cover types of interest to the analyst.

The two types of classes described above, information classes and spectral classes, may not exactly correspond to each other. For instance, two information classes, corn and soya beans, may look alike spectrally. We would say that the two classes are not separable spectrally. At certain times of the growing season corn and soya beans are not spectrally distinct while at other times they are. On the other hand a single information class may be composed of two spectral classes. Differences in planting dates or seed variety might result in the information class "corn" being reflectance differences of tasseled and un tasseled corn. To be useful, a class must be of informational value and be separable from other classes in the data.

