

CASE FILE COPY

A COMPARISON OF TWO APPROACHES FOR CATEGORY IDENTIFICATION AND CLASSIFICATION ANALYSIS FROM AN AGRICULTURAL SCENE

by

J. A. Schell

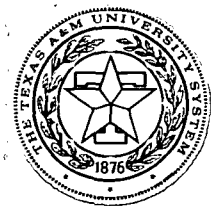
Presented At The University of Tennessee
Space Institute Conference On
Earth Resources Observations and Information
Analysis Systems

March 1972

supported by
National Aeronautics and Space Administration
Grant NsG 239-62



**TEXAS A&M UNIVERSITY
REMOTE SENSING CENTER**
COLLEGE STATION, TEXAS



Technical Report RSC-36

A COMPARISON OF TWO APPROACHES FOR
CATEGORY IDENTIFICATION AND CLASSIFICATION
ANALYSIS FROM AN AGRICULTURAL SCENE

by

J. A. Schell

ABSTRACT - Supervised and unsupervised classification modes are discussed in light of the multidisciplinary, high data rate requirements of the ERTS satellites soon to be launched. Inadequacies of each system in light of these requirements are noted and a compromise solution to the data classification system is proposed. An example of results obtained with an implementation of this system are shown and compared with results from a supervised classification scheme.

INTRODUCTION

In the field of remote sensing data analysis, considerable emphasis has been placed on the development of low cost, operational techniques for the identification of specific data classes (e.g. crop and soil type, water, etc.) from the remote multispectral measurements of the terrain. Most commonly reported is the identification of crop type from measurements of visible and infrared light reflected from a solar illuminated agricultural scene. These measurements are obtained from an imaging

multispectral scanner or densitized multiband photography. Thus the data set consists of sequential resolution elements, each a data vector of integrated reflectance values over the spectral resolution bands of the instrument.

There are several specific criteria which determine the relative worth of the analysis techniques developed, however the weighting of each item is somewhat qualitative. The developed classification system should require a minimal amount of computer processing and it should be mostly automatic. The percent correct recognition should be high, or equivalently, the measurement vectors associated with one class should be separable from the measurement vectors of all other classes. Finally, the classes into which the data are separated should be relevant to the user of the output. Thus the classes assigned to the data for the generation of soils information would be generally different than those identified with a crop survey. In a sense the latter points, separability and relevancy, are independent of one another in that the clustering of the data does not assure that the clusters are related to classes of interest and similarly the selection of relevant classes does not assure separability. It is this dilemma which has generated much discussion among investigators in this area.

Currently the goal is to provide a classification system which will adequately handle the data stream from the Earth Resources Technology Satellites. Thus the system must be computationally efficient, service a number of users and provide data for multidisciplinary interpretation. The system must also accommodate the wide variation in the data from the extensive areal coverage of the satellite, since the reflectance characteristics, particularly for vegetative types, are not fixed over such an area but change rapidly with climatic, soil, and temporal variations.

In this paper the two general classification philosophies which govern quasi-operational classification systems are examined in light of the new data products requirements of the multidisciplinary user community from ERTS multispectral data. It is shown that these two philosophies are inconsistent in part with the total system requirement and a compromise classification philosophy is proposed which suggests a solution to some of the problems inadequately addressed. An example is given to demonstrate the feasibility of a classification system utilizing this philosophy and the results are compared to those obtained through the utilization of a supervised classification scheme.

SUPERVISED CLASSIFICATION

The first of two general approaches taken in the classification analysis of multispectral data has been developed at the University of Michigan Willow Run Laboratories and at the Purdue University Laboratory for Applications of Remote Sensing [3], [7], [19]. This general approach of supervised classification is characterized by the selection of subsets from the data based upon knowledge of the classes represented. These data subsets are used to compute the signatures of the categories into which the data is to be classified. In computing the signatures, the classifier is trained, that is, the parametric values needed for the classification algorithm are computed for the categories of interest. The complexity of this and associated operations is dependent upon the initial quality of the data and whether preprocessing techniques and corrections determined by calibration values are required [9]. Generally included in this process is also a reduction and optimization of the data to assure a maximum average separability of classes, utilizing as few of the spectral measurements as necessary for each data resolution cell [4].

Although many classification algorithms have been used in the processing of multispectral data, it has been reported that the maximum-likelihood, or Bayes classifier, provides acceptable classification accuracy and processing time requirements. This classifier requires the knowledge of statistical parameters estimated from the training subsets and assumes that the data from each subset has been generated by a separate unimodal stochastic process. By examining the training subsets, the statistical parameters for the process distribution function, generally assumed to be Gaussian, are determined. The implementation of the classifier algorithm computes the value of the conditional density function for each data category, evaluated at the data point to be classified, and infers the membership of the data point to that category for which the conditional density function is largest.

The maximum likelihood classification system has been semi-operational for a period of time and has proven to be effective in an operational sense for certain classes of data. This type of operational system has undergone considerable metamorphosis since its inception to improve its recognition performance and on the whole has become quite sophisticated.

A major criticism of the operational supervised classification system has been the ambiguity introduced by the selection of classes and subsequent assumption that the data from these classes was separable. In the evolution of the system, criteria have been developed to optimize the selection of reflectance bands for processing of the data into the specific classes desired and operational procedures have been specified to assure that training samples are obtained from major groupings of the data. A second criticism relating to the amount of human intervention has been circumvented by the definition and incorporation of additional decisions into the software system. A third criticism of this approach to classification analysis has not been satisfactorily resolved. This involves the inability of the system to adequately cope with the high degree of variability within specific categories over extensive areal coverage, and the wide variety of categories which exist in a large area. Several approaches have been taken to alleviate this inadequacy. Among those attempted have been the definition of a greater number of categories, the correction of classifier parameters by an adaptive process and through the improvement of calibration of the measurement systems to remove as much variability

from the total data set as is unrelated to the categories determined [6], [8], [9]. Also under investigation have been the incorporation of temporal models of vegetative systems to correct for this variability in areal data. Each of these techniques has improved the recognition capability of this classifier system but the system is still lacking as far as "universal classification" is concerned. It must be noted however that this ideal has not been shown to be achievable, or even completely desirable, and that each innovation increases the amount of compute time required.

The classification systems developed around the maximum likelihood classifier with the improvements mentioned have proven to be valuable in the classification of multispectral data into specific categories for regional measurements. However, their usefulness have not been completely evaluated for interregional measurements, and for terrain where specific categories cannot be prelocated or where categories are mixed. The applicability of this processing system to interregional extensive data sets does not follow directly from the success achieved by this system in the past as extensive modification of the system to improve its applicability to interregional areas could change the com-

pute time appreciably. Thus the advantage held by the system for area limited data sets may be lost in its generalization to include greater numbers of classes and effective classification over larger areas.

UNSUPERVISED CLASSIFICATION

A second general approach to the problem of classification of multispectral data emphasizes the determination of categories which are separable. Through an iterative or hierarchial procedure, the measurements are grouped both in the measurement space and spatially to form accumulations of data clusters in which members of a cluster are close by some measure of distance or similarity to other members of the cluster and are distant or dissimilar from members of other clusters [1], [2], [5]. Controlling cluster parameters determines the minimum cluster size and the maximum number of categories determined from the clusters. The general difference between most cluster techniques is in the measure used for the clustering. A semi-operational classification analysis system has been implemented by the University of Kansas, Center for Research, Inc.

Emphasis in the application of cluster techniques has been placed on the definition of separable categories in the measurement space, and the relationship of these categories to classes of interest in a specific application is left to chance. Optimization of the clustering procedure tends to obtain efficiency in obtaining separated accumulations, rather than efficiency in the separation of desirable classes. Also inherent in these analyses is an increase in the computer time required to establish the clusters and this increase would generally outweigh any relative advantage in improved classification for all but the simplest of distance algorithms for an operational multidiscipline system.

The majority of clustering algorithms reported require a knowledge of the total data set before the clustering procedure is initiated, since these procedures require a knowledge of each point in the set before the clustering algorithms are applied. Processing time improvements are gained by specific ordering in the clustering procedure, comparing points to the most probable or largest cluster, and grouping contiguous points initially before determining the appropriate cluster.

It is apparent that as the data set gets larger, the clustering algorithms, comparing points throughout

the data set, require even greater processing times and the compute time increases at a geometric rate. It has been reported that the total number of clusters tends to a limit over data sets of moderate length, but it is unreasonable to assume that these clusters would be applicable over extensive interregional data sets. Even if it is expected that the total number of clusters would increase at a low rate, the comparison of new measurements individually with the large number of existing clusters would further degrade the overall classification system performance in terms of compute time required. The cluster approach to data classification is in its present configuration thus restricted to moderate regional data sets in the same manner as the maximum likelihood classification systems.

The quantity and variety of identifiable information groups resulting from the application of cluster techniques is useful in the investigation of data set information content and provides considerable leverage in the understanding of the types of data produced by a multispectral system. However, it is unlikely that such an amount of information would be required in an operational system. Too, optimization of the cluster system tends to enhance cluster separation which may

not necessarily benefit the separation of the classes of interest. Thus the cluster classification systems may be too finely tuned to the data to solve the clustering problems of a specific discipline in an operational mode.

CLASSIFICATION PHILOSOPHIES

Thus both the clustering and maximum likelihood classification techniques perform well on local data sets in an operational sense. What one procedure lacks in accuracy it gains in computer processing time. It is also reasonable that the maximum likelihood approach would have the edge in an operational sense, since the system is optimized for the discrimination of specific categories of interest. Experience has shown that this particular system does this extremely well for local, regional areas.

The dilemma remains, what would constitute a reasonable analysis philosophy for the classification of multispectral data in a low cost operational system with interregional data inputs? It seems clear that for the quality of existing multispectral sensors, the use of cataloged signature characteristics for data classification would be unresultful for all but the most general classes. It also seems clear that the "universal clas-

sification system" belongs to a distant Utopia and that a now-operational system can achieve its goals only through a regionally calibrated succession of classification analyses. Thus both the supervised and unsupervised classification techniques would be able to perform satisfactorily for specific applications. However, it appears that both techniques are too restrictive to produce a data product which would be useful to a multidisciplinary community.

A compromise between the two systems seems in order. As has been reported, clusters exist in both the spatial associations of the data and in spectral associations. The maximum likelihood classifier relies almost exclusively on spatial association of the data to compute training parameters which in a sense defines "spectral clusters". Stochastic analysis techniques are then applied to associate additional spatial coordinates with these "clusters" and the training sets are thus extrapolated to the total data set. In the measurement space cluster algorithms, accumulation of spectral values are noted and clusters generated to define these accumulations. Known categories of interest are related to these accumulations through their spatial coordinates and the membership to the category to other elements

of the cluster is inferred. Thus the differences between the processing philosophies of the two classification systems are generally the domain in which the clustering occurs, and the forcing of clusters in the supervised techniques.

It has been well established that elements of most categories of interest are in spatial proximity to other elements of their categories and that this proximity of elements together with spectral similarity is sufficient to define a meaningful spatial grouping. This spatial clustering, however, is insufficient in the linking of similar categories which are disjoint spatially. Therefore spectral similarity must be used to join the spatial clusters into meaningful data subsets or categories. Thus an operational, multidisciplinary classification philosophy is clear. In those cases where meaningful categories are represented by contiguous spatial elements, it is these elements which must form the initial data groupings or clusters. Since undefined boundaries exist between sequences of differing category, the additional information required for separation must be provided by the spectral similarity between elements of a data category and the dissimilarity of other data subsets. It is also this

spectral similarity that would be used in the association of spatially disjoint subsets of particular categories. Utilization of this philosophy of classification essentially provides an areal compression of the multivariate spectral information into spatial clusters with similar spectral characteristics. No attempt would be made to specifically limit the spatial clusters to specific problem defined categories of interest or to artificial data groupings. Specific data categories should be readily identifiable from the cluster groups by association with known spatial category distributions. The spectral similarity weighting also provides a degree of spectral separation between the spatial clusters. The spectral linking of disjoint spatial clusters allows the extrapolation of training category assignment throughout the data set.

Thus, the "spatial-spectral clustering" classification philosophy provides a compromise solution to the multispectral data classification dilemma. This approach incorporates the "relevancy of classes" of the supervised mode with the speed and unbiased of the one-pass unsupervised methods. Since two dimensional spatial proximity is emphasized, extremes of an amoebic smear in the measurement space may be divided if the spatial clusters

show that the accumulations are spatially separated. Yet the arbitrary definition of these spatial clusters is avoided. As with both other methods, the classification categories must be related to subject classes, however, in this spatial clustering method, additional classes are identified only as they occur in the data set rather than previous to the analysis as would be necessary in a supervised classification mode.

CLASSIFIER IMPLEMENTATION

An unsupervised classification procedure was implemented to demonstrate the feasibility of the modified classification philosophy. The procedure is in a sense similar to that reported by Nagy et al. [11], and is a one-pass clustering procedure. The underlying philosophy is different from the Nagy procedure in that spatial clustering is weighted more strongly initially than spectral clustering and the two dimensional spatial correlation of similar resolution elements is used for data compression rather than the one-dimensional strip formation reported. This procedure is also innovative in that the linking of spatial clusters is delayed until either the clusters have been determined for the total data set or a fixed

number of clusters has been obtained. Thus the comparison of new local clusters with old clusters is delayed until the maximum local data compression has been obtained resulting in fewer operations and improved computer efficiency.

The distance function used for the data clustering is the l_∞ metric $d = \max |X_i - X_j|$ where the distance in the measurement space is the maximum of the absolute values of the differences between the spectral components. It was found that this distance measure provided an increased sensitivity to spectral change, suppressed the smearing of categories in the measurement space, and its success could be more readily predicted from individual channel data characteristics than other measures averaging all the channels. This measure seemed also to be more characteristic of the type of variation expected as classes changed, although unfortunately it was also sensitive to spurious noise spikes.

Spatial clustering was determined by computing the spectral cluster center and by adding points to the cluster whose distance from the cluster center was less than a threshold θ_T determined *a priori*. In this way sequential data points were linked in a spatial cluster. When a data point is encountered whose distance from cluster was larger than the threshold value, the distance

to the cluster to which the adjacent cell in the previous line belongs was computed. If this distance was less than the threshold, this previous cluster was enlarged until another point with distance greater than the threshold was encountered. If the adjacent clusters are at a greater distance than the distance threshold, a new cluster was initiated. Thus locally compressed clusters were generated with only comparison to spatially adjacent clusters for the linking of clusters.

Since the bookkeeping function for determining when a spatial cluster is complete for a two dimensional irregular cluster would be quite complex, the linking of spatially disjoint, spectrally similar clusters was deferred until after the data set was completely processed or until a predetermined number of clusters were generated. At this point all of the generated clusters were compared and all clusters closer than θ_L were assigned to the same cluster. The linking parameter θ_L was considerably smaller than the original parameter θ_c to prevent linking of separate spectral clusters by a smear of points or clusters between them. Additionally a threshold was set to exclude clusters with very few members from consideration. This additional exclusion would eliminate, for the most part, noise spikes and would suppress the linking of dissimilar classes by clusters along their spatial boundaries.

DATA AND RESULTS

The data set used for the feasibility demonstration was obtained over the Weslaco, Texas test site by the University of Michigan M-5 multispectral scanner in May, 1966. The mission was flown at an altitude of 2000 ft. and the digitized data tape was obtained from the Purdue University Laboratory for Applications of Remote Sensing through the cooperation of the USDA-ARS facility at Weslaco. Specific channel assignments of the M-5 scanner are shown in Table 1 and a photograph of the area analyzed is shown in Figure 1.

Preprocessing of the data was accomplished by the division of each element of a data cell by the sum of the elements in the cell. Calibration corrections were also made on the data. To determine initial parameter values for the threshold parameter θ_T and the linking parameter θ_L , histograms of the data values for each of the twelve channels were constructed together with histograms of the absolute magnitude differences between similar channels for a spatially adjacent resolution elements. These histograms are shown in Figures 2 and 3.

As can be seen from the histograms, the majority of points fall less than 4 units in each spectral band away from their neighbors. A lesser number fall greater than 15-20 units away from their neighbors. Since it is expected that near neighbors most probably belong to the same class, it is seen from the histograms that the variation in a class should on an average be less than 4 units from the mean while variation between classes should be on an order greater than 20 units. These distributions are fairly consistent for each of the twelve channels. The two threshold values were thus selected based on the histogram data. The threshold parameter θ_T was set equal to 18, and the linking parameter θ_L was set equal to 4. These parameter values provided the best results in the classifications although other parameter values were tried. Additionally the cluster size threshold was arbitrarily set at 5 units. Clusters with less than 5 members were ignored.

In Figure 4 is shown supervised classification results from approximately a 400 line section of the Weslaco data set. These results were obtained from the application of a maximum likelihood classifier to the data which was trained on twelve different classes within the data set. Eight of these classes are represented in this particular segment of the data.

These results were obtained from a previous study where training samples were carefully selected based upon histogram distribution and ground truth documentation [12]. A subset of the twelve measurements was selected for classification from among the channels providing the best separation of the classes. The percentage correct recognition tabulation (Table 2) is based on a point by point comparison rather than per field and is presented here merely as an indication of how well the classifier worked since the tabulation is based on the total Weslaco data set rather than solely on the segment treated here.

In Figure 5 is shown the classification results from the unsupervised classification system developed from the spatial-spectral clustering philosophy. These results are presented to demonstrate the feasibility, such an approach in the classification of large data sets with the elements of classes being for the most part spatially contiguous. This classification system is not completely operational as more work is required in the development of software to calibrate the identified clusters and to further restrict the total number of clusters to be considered over large data sets to the capacity of the computing machinery.

What is readily apparent from these results is that for the most part, points in individual fields are assigned to generally a single class. That class is separable from other classes if the new signature is outside the thresholds specified for similar sets. Also classes with a high degree of similarity are identified as belonging to the same category. In a larger data set it is expected that this would be more evident. It was also noted (see Table 3) that not only are clusters generated by crop variability, but are also expanded by variability in the amount of ground cover and in basic crop validity. By a proper selection of subsets of the dimensions processed here, these variabilities may be emphasized or subdued to enhance the resultant data product according to a specific user. But most significant, this processing system has reduced this data set to basic clusters directly related to spatial accumulations. These clusters may be further enhanced or may be calibrated with known ground data and class distribution inferred.

Processing time for this unoptimized Fortran software package is approximately 4 ms per data point on an IBM 360/65, and includes processing of all 12 data channels. This is considerably less than that required

by the maximum likelihood processor and can conceivably be reduced still further through the application of programming refinements and the processing of fewer data channels.

CONCLUSION

Two basic classification systems, the supervised technique and the unsupervised techniques, were found to be applicable to the problem of classifying data with regional data sets. However, each has specific disadvantages when applied to large interregional data sets expected from the ERTS satellites and are limited by their specific design philosophy. A compromise philosophy is advanced and an example of its implementation is shown. This illustration demonstrates the feasibility of a spatial-spectral cluster system and compares results obtained by this system with those obtained from a supervised classifier for an agricultural scene.

REFERENCES

- [1] G. H. Ball, "Data Analysis in the Social Sciences: What About the Details?" 1965 Proc. Fall Joint Computer Conference.
- [2] R. E. Bonner, "On Some Clustering Techniques," IBM J. of R & O, Vol. 8, pp. 22-23, January 1964.
- [3] K. S. Fu, D. A. Landgrebe, and T. L. Phillips, "Information Processing of Remotely Sensed Agricultural Data," Proc. of the IEEE, Vol. 57, No. 4, pp. 640-653, April 1969.
- [4] K. S. Fu, P. J. Min, and T. J. Li, "Feather Selection in Pattern Recognition," IEEE Trans. on System Science and Cybernetics, Vol. SSC-6, No. 1, pp. 33-39, January 1970.
- [5] R. M. Haralick and G. L. Kelley, "Pattern Recognition with Measurement Space and Spatial Clustering for Multiple Images," Proc. IEEE, Vol. 57, No. 4, pp. 654-665, April 1969.
- [6] R. M. Hoffer and F. E. Goodrich, "Variables in Automatic Classification Over Extended Remote Sensing Test Sites," Proc. 7th Symp. on Remote Sensing, Vol. 3, pp. 1967-1974.
- [7] M. R. Holter and W. L. Walfe, "Optical Mechanical Scanning Techniques," Proc. IRE, Vol. 47, No. 9, pp. 1546-1550, 1959.
- [8] F. J. Kriegler, "Implicit Determination of Multispectral Scanner Data Variation over Extended Areas," Proc. 7th Symp. on Remote Sensing, Vol. 1, pp. 759-778.
- [9] F. J. Kriegler, W. A. Malila, R. F. Nalepha, and W. Richardson, "Preprocessing Transformation and Their Effects on Multispectral Recognition," Proc. Sixth Int. Symp. on Remote Sensing of Environment, p. 97, October 1969.

- [10] R. E. Marshall and F. J. Kriegler, "Multispectral Recognition Techniques Present Status and a Planned Hybrid Recognition System," Proc. of the 1970 IEEE Symp. on Adaptive Processes Decision and Control, p. XI.2.1., December 1970.
- [11] G. Nagy, G. Shelton, J. Tolabe, "Procedural Question in Signature Analysis," Proc. 7th Int. Symp. on Remote Sensing, Vol. 2, pp. 1387-1401.
- [12] D. A. White, J. W. Rouse, Jr., and J. A. Schell, "Analysis of Simulated Multispectral Data from Earth Resources Satellites," Proc. IEEE 3rd Int. Geoscience Electronics Symp., 1971.

TABLE 1
SCANNER BANDS

<u>Channel</u>	<u>Spectral Response (Microns)</u>
1	0.40 - 0.44
2	0.44 - 0.46
3	0.46 - 0.48
4	0.48 - 0.50
5	0.50 - 0.52
6	0.52 - 0.55
7	0.55 - 0.58
8	0.58 - 0.62
9	0.62 - 0.66
10	0.66 - 0.72
11	0.72 - 0.80
12	0.80 - 1.00

TABLE 2
SUPERVISED CLASSIFICATION RESULTS
(Weslaco - May, 1966)

<u>Class</u>	<u>Training Samples</u>	<u>Test Samples</u>
Water (1)	93.6%	72.2%
Sorghum (2, 5, 8, 9)	62.7%	76.9%
Cotton (4, 7, A)	82.3%	66.8%
Fallow (6, C)	67.7%	42.4%
Corn (3)	83.8%	-
Cabbage (B)	82.2%	-

TABLE 3
MAJOR CLUSTERS UNSUPERVISED CLASSIFICATION
(Weslaco - May, 1966)

<u>Symbol</u>	<u>Channel</u>											
	<u>1</u>	<u>2</u>	<u>3</u>	<u>4</u>	<u>5</u>	<u>6</u>	<u>7</u>	<u>8</u>	<u>9</u>	<u>10</u>	<u>11</u>	<u>12</u>
0	188	164	188	140	124	85	70	97	80	48	20	17
1	178	158	182	139	122	86	70	100	84	50	21	18
2	191	165	189	141	125	85	70	98	80	48	17	15
3	181	156	176	131	122	95	75	92	70	49	38	31
4	169	147	168	125	119	99	78	91	67	51	49	41
5	180	157	172	129	120	92	76	91	69	52	45	39
6	173	153	173	129	121	96	76	94	71	50	51	35
7	177	153	165	122	116	95	76	83	58	52	63	56
8	174	155	180	139	123	86	71	100	84	50	22	19
S	140	132	168	137	144	120	97	114	76	35	8	4



Figure 1. Nonconcurrent Aerial Photograph of Test Area.

Weslaco - 1966

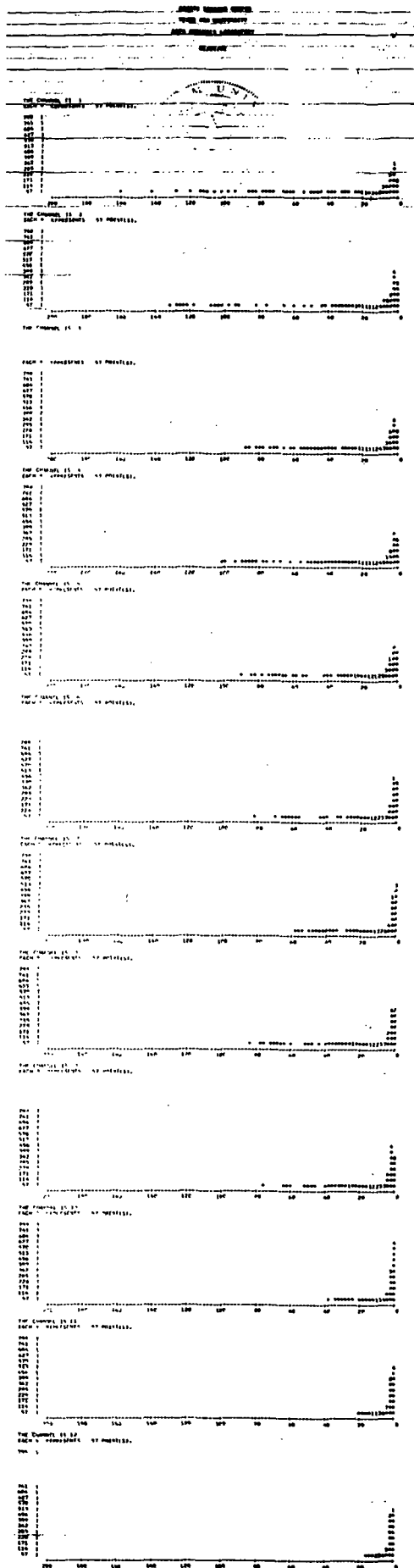


Figure 2. Data Histograms for each Channel

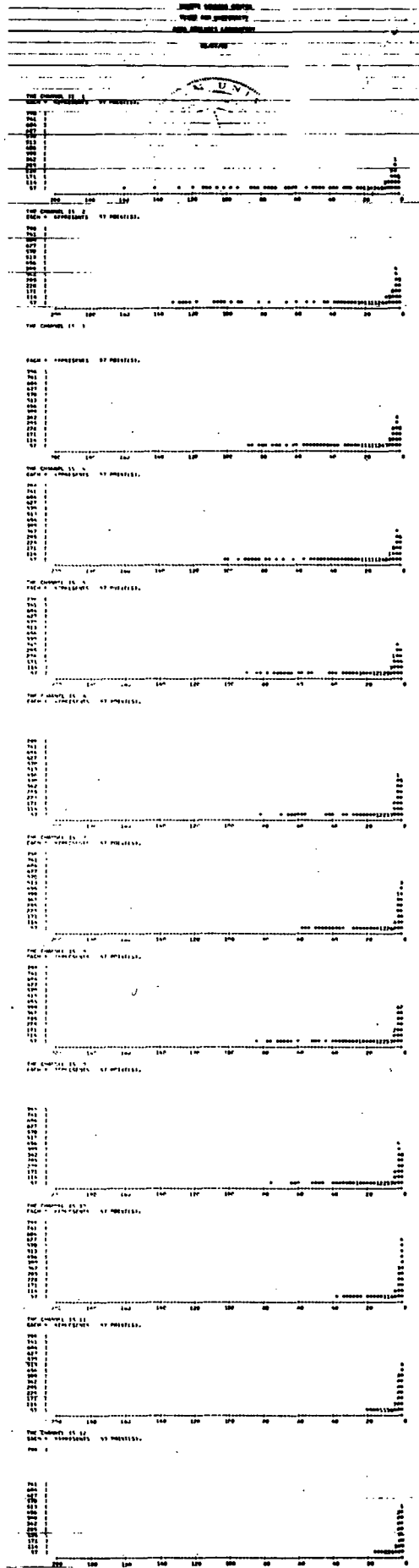


Figure 3. Histogram of Absolute Deviation
Between Spatially Adjacent Point

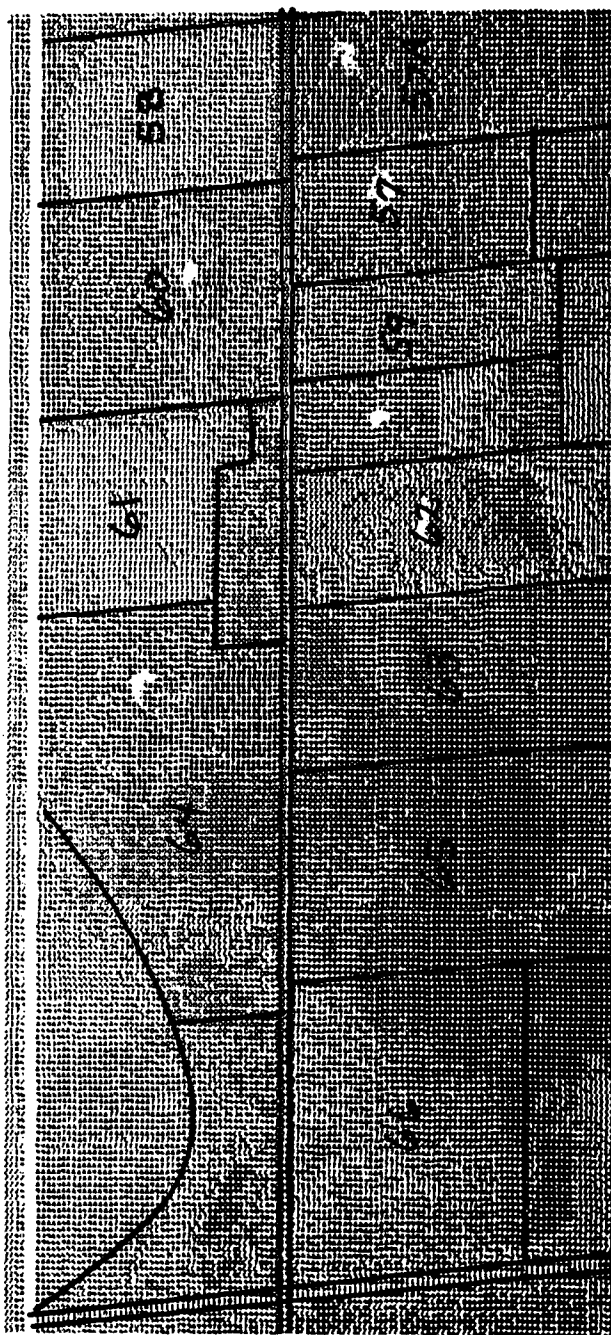


FIGURE 4

Supervised Classification Results

(Weslaco - May, 1965)

Legend

- | | |
|---------------------|---------------------|
| 0 - Unclassified | 6 - Fallow |
| 1 - Water | 7 - Cotton (70%) |
| 2 - Sorghum (25%) | 8 - Sorghum (30%) |
| 3 - Corn | 9 - Sorghum (30%) |
| 4 - Cotton (17-38%) | A - Cotton (50-63%) |
| 5 - Sorghum (96%) | B - Cabbage |
| | C - Weeds |

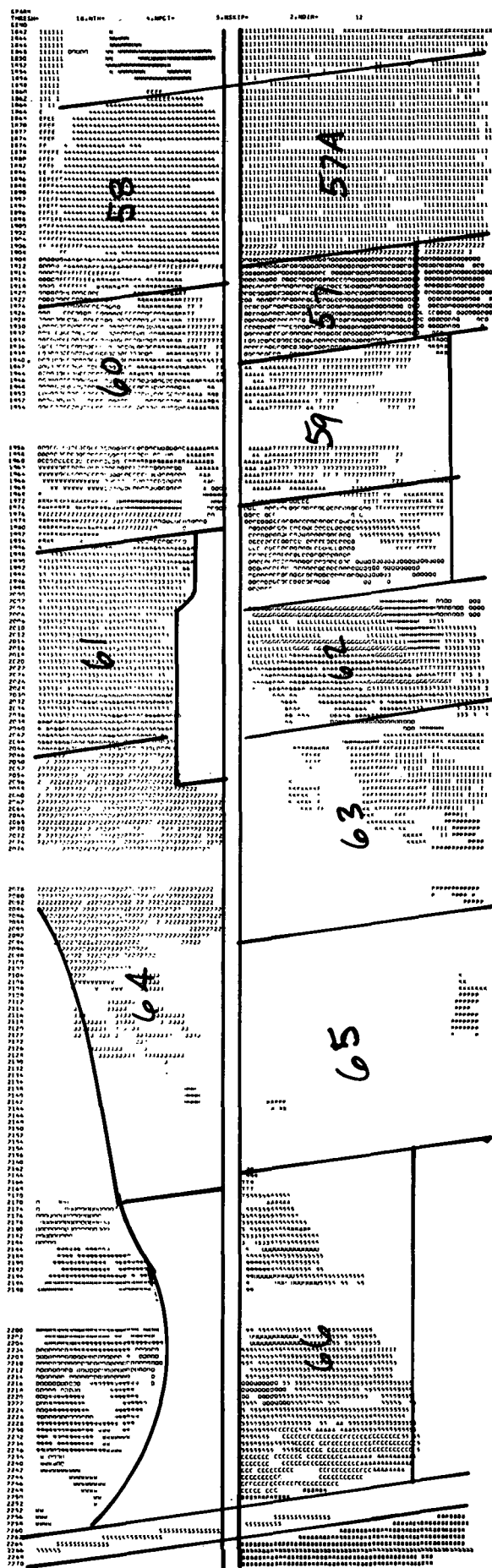


FIGURE 5

Unsupervised Classification Results

(Weslaco - May, 1966)

The REMOTE SENSING CENTER was established by authority of the Board of Directors of the Texas A&M University System on February 27, 1968. The CENTER is a consortium of four colleges of the University; Agriculture, Engineering, Geosciences, and Science. This unique organization concentrates on the development and utilization of remote sensing techniques and technology for a broad range of applications to the betterment of mankind.

