

CROP AREA ESTIMATES USING ERS-1 SAR DATA

Gerhard Smiatek

Department of Remote Sensing and Digital Image Processing

Institute of Navigation, Stuttgart University

Keplerstr. 11 , 70174 Stuttgart

Commission VII, WG 2

KEY WORDS: Remote_Sensing, Agriculture, Land_Use, Statistics, GIS, SAR, Accuracy

ABSTRACT:

Misclassified numerical results of a multispectral classification can be corrected using the double sampling scheme. For a small sample both the true classification, i.e. ground truth, and the fallible classification from the classified imagery have to be provided. The comparison of both yields information on the misclassification errors. This information is used to correct for the bias introduced by the multispectral classification. The article describes the sampling method.

KURZFASSUNG:

Die Korrektur der numerischen Ergebnisse einer multispektralen Klassifizierung kann mit der Methode der doppelten Stichprobe erfolgen. Hierbei wird für eine kleine Stichprobe die wahre Klassifizierung bereitgestellt. Aus dem Vergleich der wahren Klassifizierung und der fehlerhaften Klassifizierung mit den Bilddaten werden statistisch signifikante Korrekturfaktoren abgeleitet. Die Kenntnis dieser Faktoren erlaubt die Korrektur des Klassifikationsergebnisses. Der Beitrag beschreibt das Verfahren.

1. INTRODUCTION

The Maximum Likelihood classifier (ML) is fallible for well-known reasons. If SAR imagery is classified, the bias can be extremely large. Several methods have been developed to assess the accuracy of classification of remotely sensed data (i.e. Congalton, 1991). There are also methods available to correct for the bias introduced by multispectral classifications (i.e. Czaplewski and Catts, 1992). One of them is the double sampling developed by Tenebein (1970) employed with remotely sensed data by Card (1982), Czaplewski and Catts (1990), Smiatek (1995) and others. This article covers some issues of the double sampling for misclassified ERS-1-SAR data in agricultural surveys.

2. CORRECTION FOR THE BIAS

For a study area the geometry of the fields was digitized from orthophotograph and stored in a geographical information system (GIS). After that the site was visited and the actual land use was mapped and introduced into the database. Three multitemporal ERS-1 GTC SAR images were provided for classification. In addition, the Lee filtering was applied to the images. The results of the multispectral classification have shown that the area of major crops is widely underestimated and that neither Lee filtering nor the majority filter improved the results.

In crop surveys the acreage of a certain crop is usually of importance. Misclassified results of the ML classification can be corrected if some information on the misclassification error is known. For the study area for each ML classified pixel the true classification from the site visit is also available. Thus, the results can be summarized in the following contingency table:

		fallible maximum likelihood classifier				
		1	2	...	k	
True classifier (GIS)	1	a_{11}	a_{12}	$\cdots$	a_{1k}	$a_{1.}$
	2	a_{21}	a_{22}	$\cdots$	a_{2k}	$a_{2.}$
	3	a_{31}	a_{32}	$\cdots$	a_{3k}	$a_{3.}$
	k	a_{k1}	a_{k2}	$\cdots$	a_{kk}	$a_{k.}$
		$a_{.1}$	$a_{.2}$	$\cdots$	$a_{.k}$	n

where k is the number of land use category (i.e. winter wheat, maize *etc.*) in the survey; a_{ij} is the number of pixels whose true category is i and whose fallible category is j .

Following the misclassification model developed by Bross (1954), Tenebein (1970) and Tenebein (1972) the result of the ML classification given as the proportion π can be considered as

$$\pi = p(1 - \theta) + q\phi \quad (1)$$

where p is the true proportion and θ and ϕ are probabilities of misclassification and q is $1 - p$. This equation applies to the binomial case with only one land use category.

In case of k mutually exclusive land use categories, (multinomial case) the result for the category j is

$$\pi_j = \sum_{i=1}^k p_i \theta_{ij} \tag{2}$$

where θ_{ij} is the probability that a unit, which belongs to the category i is classified to the category j . The bias in the estimates $\pi_j - p_j$ is corrected by applying the following equation

$$\mathbf{P} = \mathbf{A} \times \mathbf{\Pi} \tag{3}$$

where $\mathbf{P}$ is $k \times 1$ column vector of the estimates p_i , $\mathbf{A}$ is $k \times k$ Matrix with the terms $n_{ij} = a_{ij}/a_{.j}$ and $\mathbf{\Pi}$ is the $k \times 1$ column vector with the results of the maximum likelihood classification. To correct the misclassified results of the ML classifier given by the vector $\mathbf{\Pi}$ the error matrix $\mathbf{A}$ must be known.

Let us consider the results of the ML classification with multitemporal ERS-SAR GTC data for five major crops. In the study area following contingency table is available from which $\mathbf{A}$ can be derived:

		fallible ML classifier						
	Crop	1	2	3	4	5	6	Σ
True	1	2751	794	907	419	120	2019	7010
	2	186	511	166	211	41	515	1630
	3	681	352	824	263	107	868	3095
	4	211	915	420	2368	460	2256	6630
	5	55	184	94	368	111	464	1276
	6	305	412	242	341	65	589	1954
Σ		4189	3168	2653	3970	904	6111	21595

where: 1 - winter wheat, 2 - winter barley, 3 - summer barley, 4 - sugar beets, 5 - maize and 6 - others

The disadvantage of the above procedure is obvious. The true classifications have to be provided for the entire area. But this data is not available. This is the aim of the image classification.

Some information can, however, be gained using sampling techniques. Suppose a small sample can be provided with true classifications. The size of that sample is denoted n . The entire area consisting of N pixels is classified using only the ML- classifier. Thus the n members of the sample are classified twice.

The results give following contingency table:

		fallible maximum likelihood classifier					
		crop	1	2	...	k	
True classifier (GIS)	1		a_{11}	a_{12}	$\cdots$	a_{1k}	$a_{1.}$
	2		a_{21}	a_{22}	$\cdots$	a_{2k}	$a_{2.}$
	3		a_{31}	a_{32}	$\cdots$	a_{3k}	$a_{3.}$
	k		a_{k1}	a_{k2}	$\cdots$	a_{kk}	$a_{k.}$
			$a_{.1}$	$a_{.2}$	$\cdots$	$a_{.k}$	n
			X_1	X_2	X_3	X_k	$N - n$

where a_{ij} is the number of pixels in the sample whose true category is i and whose fallible category is j . $X = (X_1, X_2, \cdots, X_k)$ is a vector of frequencies, where X_j is the number of pixels in the image with N units whose fallible category is j .

After Tenebein (1972), estimates of p and the misclassification probabilities θ can be derived as follows

$$\hat{p}_i = \sum_{j=1}^k a_{ij}(X_j + a_{.j})/(a_{.j}N) \tag{4}$$

$$\hat{\theta}_{ij} = (X_j + a_{.j})a_{ij}/(a_{.j}Np_i) \tag{5}$$

The results of the maximum likelihood classification of the N - members of the image $(X_j + a_{.j})$ are corrected by multiplying with the ratios $a_{ij}/a_{.j}$ and summing over j . In the case of k land use categories, the estimates are derived as given in equation 3.

In addition a coefficient of reliability K is defined. It measures the strength of the relationship between the true measurements and maximum likelihood classifier for each category.

$$K_i = p_i(\sum_{j=1}^k \theta_{ij}^2/\pi_j - 1)/q_i \tag{6}$$

The variance of p_i is

$$V(\hat{p}_i) = \frac{p_i q_k}{n}(1 - K_i) + \frac{p_i q_k}{N}K_i \tag{7}$$

If $K = 0$, the true classification is the only way to obtain reliable estimates for the fraction of considered crop area. If $K = 1$, the ML classifier is not fallible, thus the true classifications are not required. In sampling optical data K is within the range of 0.3 - 0.8 (Smiatek, 1993).

It is necessary, however, to accept the sampling error as the error probabilities are estimated from a sample. But, the required sample size can be estimated according to the desired accuracy criteria. This is the great advantage of the sampling procedure proposed here.

In the study area a random sample consisting of 1454 pixels yields

$$\begin{bmatrix} \hat{p}_1 \\ \hat{p}_2 \\ \hat{p}_3 \\ \hat{p}_4 \\ \hat{p}_5 \\ \hat{p}_6 \end{bmatrix} = \begin{bmatrix} 0.647 & 0.265 & 0.364 & 0.096 & 0.183 & 0.310 \\ 0.052 & 0.122 & 0.068 & 0.04 & 0.083 & 0.092 \\ 0.162 & 0.132 & 0.296 & 0.079 & 0.117 & 0.117 \\ 0.07 & 0.291 & 0.148 & 0.581 & 0.383 & 0.315 \\ 0.011 & 0.032 & 0.04 & 0.119 & 0.150 & 0.075 \\ 0.059 & 0.159 & 0.085 & 0.083 & 0.083 & 0.092 \end{bmatrix} \quad (8)$$

$$\begin{bmatrix} 0.1940 \\ 0.1467 \\ 0.1229 \\ 0.1838 \\ 0.0419 \\ 0.3110 \end{bmatrix}$$

and the estimated proportions are

$$\hat{P} = \begin{bmatrix} \hat{p}_1 \\ \hat{p}_2 \\ \hat{p}_3 \\ \hat{p}_4 \\ \hat{p}_5 \\ \hat{p}_6 \end{bmatrix} = \begin{bmatrix} 0.32399 \\ 0.07340 \\ 0.1401 \\ 0.3181 \\ 0.0614 \\ 0.0903 \end{bmatrix}$$

The estimates deviate from the true proportions known from the field survey only by a small margin. The reason for is the large size of the sample and sampling with the sampling unit pixel. But, the coefficient of reliability K is very low. In case of ML classification of the raw SAR data both classifiers are not correlated. The sampling accuracy depends almost entirely on the true classifications. As K approaches 0 the variance depends only on the size of the sample n (Equation 7). If K rises, the left term of Equation 7 becomes smaller with $(1 - K)$. But the right term must be added. This term is, however, very small because of large N . Thus, some expensive true observations are replaced by a large number of misclassified observations.

Larger K values are found in the results of the majority filtering. In that case the results of the ML classification of the entire imagery improve the accuracy of the estimates from the small sample with true classifications. But the majority filtering rises the entire cost of the sampling.

In the sampling scheme described above a full image classification was used. The full frame classification can be replaced by a sample with N pixels, where $N \gg n$. In that case the n members of the sample are called subsample. With a given coefficient of reliability K and the probability of misclassification θ , both n and

N can be optimized for given accuracy requirements. Tenebein (1972) describes the procedure.

3. CONCLUSIONS

From the results of this study can be concluded, that the double sampling has the capability to correct misclassification errors of the ML classification. However, the coefficient of reliability is with ERS-1-SAR-data very low. This means that a comparatively large and expensive sample with true classifications is needed to get sufficient accurate estimates of the error probability matrix and other parameters. The results apply to the sampling unit pixel. The sampling performance can be improved by a carefully trained classifier, use of systematical sampling and exclusion of non-agricultural areas from the sampling.

The author thanks ESA for providing the ERS-1 SAR imagery

References

Bross, I. (1954): Misclassification in 2*2 tables. *Biometrics* (10), 478-486.

Card, D. H. (1982): Using known map categorical marginal frequencies to improve estimates of thematical map accuracy. *Photogrammetric Engineering and Remote Sensing* 48, 431-439.

Congalton, R. G. (1991): A reiew of assessing the accuracy of classifications of remotely sensed data. *Remote Sensing of Environment* 37, 35-46.

Czaplewski, R. L. and Catts, G. P. (1990): Calibrating area estimates for classification error using confusion matrix. *Technical Papers 1990 ACSM-ASPRS Annual Convention* 4, 431-440.

Czaplewski, R. L. and Catts, G. P. (1992): Calibration of remotely sensed proportion or area estimates for misclassification error. *Remote Sensing of Environment* 39, 29-43.

Smiatek, G. (1993): *Erfassung der Flächennutzung mit Hilfe von LANDSAT/TM-Stichproben*. Dissertation Institut für Navigation der Universität Stuttgart Stuttgart. 156 S.

Smiatek, G. (1995): Sampling thematic mapper inagery for land use data. *Remote Sensing of Environment* 52, 116-121.

Tenebein, A. (1970): A double sampling scheme for estimating from binomial data with misclassifications. *JASA* 65(331), 1350-1361.

Tenebein, A. (1972): A double sampling scheme for estimating from misclassified multinomial data with applications to sampling inspection. *Technometrics* 14, 187-202.