



Probabilistische segmentatie en fuzzy classificatie van natuurlijke vegetatie in hyperspectrale beelden

Scriptie voor de opleiding Geodesie
Technische Universiteit Delft

mei 2003

Jochem Lesparre

Voorwoord

Deze scriptie is de schriftelijke verslaglegging van mijn afstudeeronderzoek voor de opleiding Geodesie aan de Technische Universiteit Delft. Het onderzoek is uitgevoerd bij de Meetkundige Dienst (MD) van Rijkswaterstaat. In de loop van 2003 zal deze dienst "Adviesdienst Geo-informatie en ICT" (AGI) gaan heten.

Helaas staat deze scriptie vol met niet-Nederlandse woorden als fuzzy, k-nearest-neighbours-classificatie en remote-sensing-beeld. Ik heb er bewust voor gekozen deze woorden niet te vertalen, omdat de tekst anders onduidelijk zou worden voor eenieder die helemaal vertrouwd is met deze termen.

Graag wil ik de Meetkundige Dienst bedanken voor de geboden onderzoeksfaciliteiten. Uiteraard ben ik ook aan enkele mensen dank verschuldigd, in de eerste plaats aan dr. ir. B.G.H. Gorte (TU-Delft), R.W.L. Jordans (MD) en prof. dr. ir. M.G. Vosselman (TU-Delft) voor hun begeleiding. In de tweede plaats wil ik ir. F. Kenselaar, drs. E.H. Kloosterman, drs. R. de Lange, dr. K. Schmidt, prof. dr. A.K. Skidmore, ir. E.M.J. Vaessen, drs. M.S. de Waal en de bewoners van Huize x-dakje bedanken voor hun adviezen of praktische hulp.

Inhoud

Voorwoord	i
Inhoud	iii
Samenvatting	v
Summary	ix
1. Inleiding	1
1.1. Aanleiding	1
1.2. Doelstelling	2
1.3. Werkwijze	3
1.4. Leeswijzer	3
2. Mixed pixels en fuzzy classificatie	5
2.1. Mixed pixels	5
2.2. Fuzzy classificatie	6
2.3. Lineaire vermenging van spectra	6
3. Beschrijving van het studiegebied en de data	9
3.1. Studiegebied	9
3.2. Beeldmateriaal	9
3.3. Veldgegevens	10
3.4. Klassenindeling	11
4. Fuzzy training voor maximum-likelihood-classificatie	15
4.1. Aanleiding	15
4.2. Nieuwe methode voor fuzzy training	16
4.3. Experiment	22
4.4. Beschouwing en aanbevelingen	25
5. Probabilistische segmentatie	27
5.1. Classificatie	27
5.2. Segmentatie	33
5.3. Beschouwing en aanbevelingen	37
6. Conditionele kans op een featurevector	41
6.1. Overtraining	41
6.2. Keuze voor een schattingsmethode	42
6.3. k-Nearest-neighbours-classificatie	43
6.4. Beschouwing en aanbevelingen	45
7. Evaluatie	47
7.1. Validatie	47
7.2. Resultaten	48
7.3. Conclusies en aanbevelingen	49
8. Conclusies en aanbevelingen	51
8.1. Fuzzy training voor maximum-likelihood-classificatie	51
8.2. Probabilistische segmentatie	51
Literatuur	53
Bijlage 1. Classificatieresultaten	55
Bijlage 2. Handleiding voor de programmatuur	69
Symbolen	77

Samenvatting

Aanleiding

De Meetkundige Dienst doet voor Rijkswaterstaat onderzoek naar semi-automatische interpretatie van remote-sensing-beelden om de productiviteit en met name de objectiviteit van de vegetatiekartering te verbeteren. Er doen zich hierbij twee problemen voor. Ten eerste lijken bij vergelijking van twee afzonderlijke classificaties door toevalligheden veranderingen op te treden die in werkelijkheid geen veranderingen zijn. Het tweede probleem is dat geleidelijke veranderingen van mengverhoudingen van klassen niet of pas laat zichtbaar zijn, omdat deze vaak geen aanleiding geven tot classificatie in een andere klasse. De oplossing voor deze twee problemen wordt gezocht in het gebruiken van fuzzy technieken.

Een ander probleem is dat automatische classificatie van natuurlijke vegetatie redelijk moeizaam is, daar de verschillende klassen vaak spectraal lastig te onderscheiden zijn. Daarom is het wenselijk gebruik te maken van hyperspectrale scanners en wordt er gezocht naar nieuwe technieken die op een slimmere manier classificeren, door bijvoorbeeld gebruik te maken van expertkennis of ruimtelijke samenhang.

Doelstelling

Doel van het afstudeerproject is te onderzoeken in hoeverre het gebruik van een fuzzy classificatie van hyperspectrale beelden met behulp van een probabilistische segmentatie [Gorte 1998] verbetering geeft ten opzichte van een crisp, pixelgewijze classificatie voor monitoring van natuurlijke vegetatie. Er wordt verwacht dat probabilistische segmentatie een verbetering geeft omdat deze methode naast de spectrale eigenschappen ook rekening houdt met de ruimtelijke kenmerken van de terrein-eenheden. De probabilistische segmentatiemethode is ontwikkeld voor multispectrale data van agrarische gebieden. Natuurlijke vegetatie groeit niet zo gestructureerd als cultuurgewassen, maar ook niet lukraak door elkaar. Daarom wordt ook voor natuurlijke vegetatie een verbetering van de classificatie door segmentatie verwacht.

Werkwijze

De gevolgde werkwijze behelst het toepassen van de probabilistische segmentatiemethode op hyperspectrale beelden van natuurlijke vegetatie. Aanvullend is een beknopt literatuur onderzoek uitgevoerd. Het operationaliseren en aanpassen van software vergt veel tijd. Daardoor ligt de nadruk van het afstudeeronderzoek op het uitbreiden van probabilistische segmentatie voor toepassing op hyperspectrale beelden van natuurlijke vegetatie en het programmeren hiervan.

Tijdens het onderzoek ontstond het idee voor een nieuwe methode voor fuzzy training voor maximum-likelihood-classificatie. Uiteindelijk is er voor k-nearest-neighbours-classificatie gekozen in plaats van een maximum-likelihood-classificatie. Daarom is de fuzzy trainingsmethode alleen beschreven en niet toegepast.

Data

Het gebruikte beeldmateriaal is ingewonnen met het HyMap-systeem (Hyperspectral Mapping System). Dit is een hyperspectrale scanner voor in een vliegtuig die 128 banden opneemt. Hiermee is de kwelder van Schiermonnikoog opgenomen. De ruimtelijke resolutie van de beelden is 3,5 meter.

Voor de training en validatie van de classificatie is gebruik gemaakt van 384 tijdens veldwerk ingewonnen veldopnamen. Deze veldopnamen zijn ongeveer half om half gesplitst in

trainingssamples en validatiesamples. Tijdens het onderzoek zijn vier verschillende klassen-indelingen voor de veldopnamen gebruikt. Het opstellen van deze indelingen is handmatig geschied, hierdoor zijn deze redelijk subjectief. Dit resulteert in klassen met een relatief grote spectrale overlap.

Fuzzy training voor maximum-likelihood-classificatie

De meeste bestaande technieken hebben pure pixels nodig voor de training. Dit maakt de training voor natuurlijke vegetatie lastig. Want door het voorkomen van veel mixed pixels is het moeilijk om voldoende pure trainingssamples te vinden. Het gebruik van bijna pure trainingssamples is suboptimaal. Hierbij wordt namelijk een deel van de geschatte spreiding van een klasse niet veroorzaakt door de natuurlijke variatie in het spectrum van die klasse, maar door de mate waarin aandelen van andere klassen aanwezig zijn. De oplossing voor een gebrek aan pure trainingssamples moet gezocht worden in het schatten van pure spectra uit mixed trainingssamples. F. Wang [1990] stelt een methode voor waarbij een gewogen gemiddelde en een gewogen empirische (co)variantie worden berekend, met de fractie van de betreffende klasse als gewicht. Deze procedure leidt echter niet tot zuivere schattingen voor de spectra van de pure klassen.

Met behulp van de vereffeningstheorie en kansmodelschatting is het wel mogelijk een fuzzy maximum-likelihood-training uit te voeren die uit mixed pixels spectra van pure klassen schat zonder systematische fout. Voordeel van fuzzy training is dat meer pixels in het beeld bruikbaar zijn voor training. Hierdoor is het mogelijk om heterogene gebieden te gebruiken voor training of om de trainingssamples op random plaatsen in te winnen. Voorwaarde voor het gebruik van deze methode voor fuzzy training is dat men beschikt over schattingen van de klassenaandelen van de mixed trainingssamples en dat de klassen in een pixel spectraal lineair vermengd zijn. Om de bruikbaarheid van fuzzy training voor maximum-likelihood-classificatie te beoordelen verdient het aanbeveling om de methode te testen op een beeld waarvoor tijdens veldwerk de klassenaandelen van fuzzy trainingssamples zijn vastgesteld.

Probabilistische segmentatie

Probabilistische segmentatie bestaat uit een integratie van beeldsegmentatie en -classificatie. De classificatie vindt plaats op grond van de formule van Bayes. De a priori kans op een klasse in deze formule wordt vaak voor alle klassen gelijkgehouden, om de classificatie niet te veel te sturen in de richting van de meest voorkomende klassen. Bij probabilistische segmentatie wordt hier juist wel gebruik van gemaakt. De methode schat deze kans lokaal uit het beeld met behulp van het resultaat van een statistische classificatiemethode. Door deze kans lokaal te schatten kan de classificatie verbeterd worden. Het is bijvoorbeeld waarschijnlijker dat een geel pixel in een geel veld tarwe dan gras is, terwijl een zelfde geel pixel in een groen veld waarschijnlijker (verdroogd) gras dan tarwe is. Hiervoor is een opdeling van het beeld in segmenten nodig, waarbij de segmenten zo veel mogelijk samenvallen met terreinobjecten zodat er slechts één of enkele klassen per segment aanwezig zijn.

De gebruikte segmentatiemethode voegt op basis van de spectrale kenmerken van drie banden aangrenzende pixels samen tot segmenten. Pixels worden samengevoegd indien de Euclidische afstand in de featurespace tussen de pixels kleiner is dan een drempelwaarde en de (co)varianties van de gecreëerde segmenten niet groter worden dan een andere drempelwaarde. Een probleem is dat het onduidelijk is welke drempelwaarde de beste segmentatie oplevert. Bovendien is voor verschillende delen van het beeld een andere drempelwaarde het beste. De oplossing hiervoor is het maken van een segmentatiepiramide door te segmenteren met oplopende drempelwaarden. Vervolgens worden uit de verschillende niveaus van de piramide segmenten geselecteerd, op basis van de klassenaandelen in de segmenten. Hiervoor moeten voor de gehele piramide de klassenaandelen per segment geschat worden, waarvoor de methode voor het schatten van de a priori kansen op de klassen gebruikt wordt. Er worden segmenten geselecteerd met zo min mogelijk klassen. Als er meerdere kandidaten zijn worden

zo groot mogelijke segmenten gekozen, omdat de schattingen van de a priori kansen voor grote segmenten nauwkeuriger zijn.

Conclusies

Voor de gebruikte beelden en bijbehorende veldgegevens geeft de k-nearest-neighbours-classificatie de best geschatte kansen als invoer voor de probabilistische segmentatie. Voor toepassing van probabilistische segmentatie op hyperspectrale beelden van natuurlijke vegetatie met een beperkt aantal trainingssamples zijn twee aanpassingen gedaan. Ten eerste is er gekozen voor een implementatie van de k-nearest-neighbours-classificatie die hyperspectrale data kan verwerken. Ten tweede is het segmentselectiecriterium zo aangepast dat de voorkeur niet langer alleen naar pure segmenten uit gaat, maar dat er gestreefd wordt naar segmenten met zo min mogelijk klassen.

Probabilistische segmentatie geeft een lichte verbetering van de overall-accuracy van de classificatie. Deze verbetering is afhankelijk van de klassenindeling. De grootste verbetering treedt op bij een indeling in 9 klassen (niveau 1), waarbij de overall-accuracy van 61,93% naar 67,43% stijgt. De indeling met de kleinste verbetering heeft 21 klassen (niveau 3). Hiervoor steeg de overall-accuracy van 40,83% naar 41,28%. Het is onduidelijk waarom de verbetering door probabilistische segmentatie slechts klein is. Dit kan aan de data liggen of aan de probabilistische segmentatiemethode. Het eerste kan veroorzaakt worden door een te grote spectrale overlap tussen de klassen, maar het kan ook liggen aan een te lage ruimtelijke resolutie van het beeld of het geringe aantal trainingssamples. Dat de segmentselectie van de probabilistische segmentatiemethode een te grote voorkeur heeft voor pure segmenten (ook als deze heel klein zijn) kan ook een oorzaak zijn.

Aanbevelingen

Op de eerste plaats zou onderzocht moeten worden wat de oorzaak is dat de verbetering door probabilistische segmentatie slechts klein is, en hoe het komt dat veel erg kleine segmenten geselecteerd worden. Door het toepassen van probabilistische segmentatie op andere data kan geverifieerd worden of de data het probleem zijn. Door gebruik te maken van een handmatig gedefinieerde segmentatie zou gecontroleerd kunnen worden of de segmentselectiemethode de oorzaak is.

Verder is het sterk aan te raden om een methode te ontwerpen waarmee een fuzzy classificatie gevalideerd kan worden.

Summary

Introduction

The Survey Department (Meetkundige Dienst) investigates semi-automated interpretation of remote sensing images for Rijkswaterstaat to improve the efficiency and especially the objectivity of its vegetation mapping. This shows two problems. First, there seem to be differences when comparing two successive classifications that are caused by chance, not by real changes. The second problem is that gradual changes in mixing proportions of two classes are detected not at all or too late, because they often give no cause for classification in another class. The solution to these two problems is expected to be found in fuzzy classification methods.

Another problem is the poor spectral distinction between classes of natural vegetation, which complicates automated classification. Therefore one desires the use of hyperspectral scanners and researches the use of new smart classification methods that for instance make use of expert knowledge or spatial structures.

Objective

The objective of this master thesis is to investigate to what extent the use of fuzzy classification by means of probabilistic segmentation [Gorte 1998] gives an improvement compared to crisp pixel-by-pixel classification for monitoring of natural vegetation. An improvement by probabilistic segmentation is expected because this method takes into account spatial characteristics of the field objects. Probabilistic segmentation is developed on multispectral data of agricultural areas. Natural vegetation does not grow that structured as cultivated crops, but it is not completely random either. Hence, an improvement of the classification by segmentation is expected for natural vegetation too.

Methodology

The used method consisted of the application of the probabilistic segmentation method on hyperspectral images of natural vegetation. Additionally, a brief literature research was executed. As the development and modification of software take a lot of time, this research is focused on the extension of probabilistic segmentation for application on hyperspectral images of natural vegetation, and programming this.

During the research the idea arose for a new method of fuzzy training for maximum likelihood classification. As it was decided to use k-nearest neighbours classification instead of maximum likelihood classification in the end, the fuzzy training method is only described, and has not been applied.

Data

The used images are acquired by the HyMap (Hyperspectral Mapping) system. This is a hyperspectral airborne scanner that records 128 bands. This scanner is used to acquire the salt marsh on the Dutch Friesian Island Schiermonnikoog with a spatial resolution of 3.5 meters.

For training and validation of the classification 384 field plots are used which were acquired during fieldwork. These have been split about fifty fifty in a set of training samples and validation samples. In this study four different divisions in classes are used. The appointment of these divisions is made manually, which makes them quite subjective. This results in classes with rather large spectral overlaps.

Fuzzy training for maximum likelihood classification

Most existing methods need pure pixels for training, which complicates training for natural vegetation. Because of the occurrence of many mixed pixels it is difficult to obtain a sufficiently large number of pure training samples. The use of almost pure samples is not advisable. These cause the estimated variance to be composed partly by the degree of occurrence of other classes, instead of the natural variation in the spectrum of the class. The solution for the lack of pure training samples is to be found in the use of mixed pixels to estimate the spectra of pure classes. F. Wang [1990] suggests a method to calculate a weighted average and (co)variance by using the fraction of a class as the weight. This method does not give an unbiased estimation of the spectra of the pure classes.

It is possible to execute an unbiased maximum likelihood training that estimates pure spectra out of mixed pixels by using the adjustment theory and probability model estimation. The advantage of fuzzy training is that more pixels in the image can be used for training, which enables the use of heterogeneous areas for training or random sampling. There are two conditions on this fuzzy training method. First, one needs to have estimations of the fractions of the classes for the mixed training samples. Secondly, the spectra of the mixed pixels should be a linear mixture of the composing classes. In order to test the utility of fuzzy training for maximum likelihood classification, it is recommended to apply the method to an image for which class fractions of the training samples have been determined during fieldwork.

Probabilistic segmentation

Probabilistic segmentation is an integration of image segmentation and classification. The classification is carried out according to Bayes formula. The prior probabilities of a class in this formula are normally kept equal for all classes, so the classification is not favoured too much towards the prevailing classes. Probabilistic segmentation, however, does do this. The method estimates this probability locally on the basis of the image using the result of a probabilistic fuzzy classification. By estimating this probability locally it's possible to improve the classification. A yellow pixel in a yellow field for instance is more likely to be wheat than grass, while the same yellow pixel in a green field is more likely to be (parched) grass than wheat. To estimate the local prior probabilities a subdivision of the image in segments is needed, in which segments coincide with field objects as much as possible. So only one or a few classes are present per segment.

The applied segmentation method merges adjacent pixels to segments on the basis of the spectral characteristics of three bands. Pixels are merged if the Euclidean distance in the feature space between the pixels is smaller than a threshold, and the (co)variances of the created segment is smaller than another threshold. It's not clear, however, which threshold gives the best segmentation. Moreover, the best threshold seems to vary for different areas in the image. The solution to this is building a segmentation pyramid by segmenting the image with increasing thresholds. Next, segments are selected from different levels of the pyramid, on the basis of the class proportions in the segments. To do this, class proportions per segment are needed for the entire pyramid, which are acquired by the estimation method for the prior probabilities. The segments with as few as possible classes are selected. If there is more than one candidate, the largest segment is selected, because the estimates of the prior probabilities are more accurate for large segments.

Conclusions

For the utilised image and accompanying field data the k-nearest neighbours classification gives the best estimated probabilities as input for probabilistic segmentation. Two modifications were carried out for applying probabilistic segmentation on hyperspectral images of natural vegetation with a limited number of training samples. First, an implementation of k-nearest neighbours classification that's able to deal with hyperspectral

data is adopted. Secondly, the segment selection rules are adapted in a way that no longer only pure segments are preferred, but segments with as least classes as possible are aimed for. Probabilistic segmentation gives a slight improvement in the overall accuracy of the classification. The rate of improvement depends on the used division in classes. The maximum improvement occurred with a division in 9 classes (level 1). The overall accuracy increased from 61.93% to 67.43%. The division with the minimum improvement has 21 classes (level 3). Here the overall accuracy increased from 40.83% to 41.28%. It is unclear why the improvements are quite small. The data or the probabilistic segmentation method could be the cause. The data can cause this by a too large spectral overlap of the classes, but also by a too low spatial resolution or by too few training samples. A too large preference for pure segments (even when these are really small) in the segment selection of the probabilistic segmentation method could also cause this.

Recommendations

In the first place the cause of the minor improvement by probabilistic segmentation should be examined, as well as the reason for the selection of many small segments. By applying probabilistic segmentation on other data it could be verified whether the data is the problem. By using a manually defined segmentation it could be checked whether the segment selection method is the cause.

Furthermore, the design of a method to validate a fuzzy classification could be strongly recommended.

1. Inleiding

1.1. Aanleiding

Eén van de taken van Rijkswaterstaat is de verdediging van de Nederlandse kust tegen de zee. In kustgebieden liggen vaak waardevolle natuurgebieden zoals kwelders. Dit zijn begroeide, zoute of brakke gebieden op de grens van land en zee (zie figuur 1.1). Sinds de middeleeuwen zijn door het bouwen van dijken, ten behoeve van veiligheid en landbouwgrond, veel kwelders verloren gegaan. De laatste decennia is echter een ontwikkeling op gang gekomen waarbij men de natuur meer vrij laat. Voorwaarde is natuurlijk dat de veiligheid van het land gewaarborgd blijft. Daarom is monitoring van de natuurlijke processen geboden, omdat die een indicatie geven voor terreinverandering. Hierdoor heeft Rijkswaterstaat ook een natuurbeheertaak.

Voor monitoring van de Nederlandse kwelders vervaardigt de Meetkundige Dienst van Rijkswaterstaat vegetatiekaarten door visuele interpretatie van false-colour-luchtfoto's. Hoewel deze methode effectief is wordt er onderzoek gedaan naar semi-automatische interpretatie van remote-sensing-beelden om de productiviteit en met name de objectiviteit van de vegetatiekartering te verbeteren [Skidmore e.a. 2001 blz. 5].

De gebruikelijke geautomatiseerde methoden voor vegetatieclassificatie zijn echter niet zo geschikt voor monitoring van kleine of geleidelijke veranderingen [Janssen 2001 blz. 16]. Er doen zich namelijk twee problemen voor. Ten eerste lijken bij vergelijking van twee afzonderlijke classificaties veranderingen op te treden die eigenlijk geen veranderingen zijn, maar waar door een net iets andere bemonstering van mixed pixels en andere toevalligheden verschillende klassen toegekend worden aan overeenkomstige pixels. Het tweede probleem is dat geleidelijke veranderingen van mengverhoudingen niet of pas laat zichtbaar zijn, omdat deze vaak geen aanleiding geven tot classificatie in een andere klasse.

De oplossing voor deze problemen wordt gezocht in het gebruiken van fuzzy technieken. Hierbij worden mengsels van de vegetatieklassen verondersteld, dit in tegenstelling tot crisp classificaties waarbij slechts één klasse per element bepaald wordt. Hoewel eerdere projecten van de Meetkundige Dienst van Rijkswaterstaat al gebruik maakten van fuzzy technieken [Droesen 1999] ontbreekt er nog een operationeel fuzzy classificatieprogramma.

Een ander probleem is dat automatische classificatie van natuurlijke vegetatie tamelijk moeizaam is, daar de verschillende klassen vaak spectraal lastig te onderscheiden zijn. Daarom is het wenselijk gebruik te maken van hyperspectrale scanners en wordt er gezocht naar nieuwe technieken die op een slimmere manier classificeren, door bijvoorbeeld gebruik te maken van expertkennis of ruimtelijke samenhang.



Figuur 1.1. Kwelder op Schiermonnikoog

1.2. Doelstelling

Doel van het afstudeerproject is te onderzoeken in hoeverre het gebruik van een fuzzy classificatie van hyperspectrale beelden met behulp van een probabilistische segmentatie [Gorte 1998] verbetering geeft ten opzichte van een crisp, pixelgewijze classificatie voor monitoring van natuurlijke vegetatie.

Er wordt verwacht dat probabilistische segmentatie een verbetering geeft omdat deze methode naast de spectrale eigenschappen ook rekening houdt met de ruimtelijke kenmerken van de terrein-eenheden. Probabilistische segmentatie gebruikt hiervoor een integratie van beeldsegmentatie en -classificatie. De segmentatie verdeelt een beeld in aaneengesloten groepen pixels. Hierbij is het de bedoeling dat deze overeenkomen met terreinobjecten, zodat er slechts één of enkele klassen per segment aanwezig zijn. Een statistische fuzzy classificatie per pixel voorziet deze segmenten van een schatting van de klassenaandelen. De integratie zorgt ervoor dat zo homogeen mogelijke segmenten gevormd worden. Vervolgens kunnen de klassenaandelen per segment gebruikt worden om de statistische fuzzy classificatie van de pixels binnen het segment te verbeteren.

De probabilistische segmentatiemethode is ontwikkeld voor multispectrale data van agrarische gebieden. Dat het onderscheiden van objecten in agrarische gebieden de classificatie verbetert, ligt voor de hand. De verschillende gewassen staan namelijk meestal op afzonderlijke akkers. Natuurlijke vegetatie groeit niet zo gestructureerd, maar ook niet lukraak door elkaar. Componenten van natuurlijke landschappen kunnen realistisch worden gemodelleerd als een mozaïek van scherp begrensde ruimtelijke objecten met gradueel variërende interne velden [Droesen 1999]. Daarom wordt ook voor natuurlijke vegetatie een verbetering van de classificatie door segmentatie verwacht. Om de gradueel variërende interne velden accuraat te beschrijven is fuzzy classificatie benodigd. Het gebruik van hyperspectrale in plaats van multispectrale data is wenselijk omdat deze extra informatie geeft.

Randvoorwaarden en uitgangspunten

Er is gekozen om gebruik te maken van HyMap-beelden van het kweldergebied op oostelijk Schiermonnikoog. Deze beelden zijn eenmalig opgenomen. Als gevolg hiervan zal het niet mogelijk zijn om uitspraken te doen over het analyseren van multi-temporele data. Daar het wel het doel van monitoring is om veranderingen in de tijd te beschrijven is dit een onderwerp voor eventueel vervolgonderzoek.

Deze zelfde HyMap-beelden zijn eerder gebruikt voor het opstellen van een expertsysteem dat bij het ITC (International Institute for Geo-Information Science and Earth Observation, Enschede) is ontwikkeld in samenwerking met de MD [Skidmore e.a. 2001]. Bij dit expertsysteem wordt er net als bij probabilistische segmentatie een fuzzy classificatie verkregen door de kansen op de verschillende klassen te berekenen. Bij het expertsysteem verbeteren extra kansen op grond van aanvullende gegevens, zoals hoogte boven zeeniveau of kalkgehalte in de grond, het resultaat. De probabilistische segmentatie maakt geen gebruik van dergelijke aanvullende informatie, maar van de ruimtelijke samenhang om extra kansen te berekenen die het resultaat te verbeteren.

Het commerciële programma eCognition (van Definiens Imaging) gebruikt ook een integratie van beeldsegmentatie en -classificatie. Er is gekozen om ondanks het bestaan van een kant-en-klaar programma als eCognition toch het nog niet uitontwikkelde programma voor probabilistische segmentatie te gebruiken. De eerste reden hiervoor is dat het onduidelijk is wat eCognition precies doet omdat dit om bedrijfstechnische redenen geheim gehouden wordt. Een tweede reden is dat eCognition de verkregen segmenten classificeert op grond van het gemiddelde spectrum van dat segment, terwijl probabilistische segmentatie de spectra van alle pixels in ogenschouw neemt voor het classificeren van een segment.

1.3. Werkwijze

De gevolgde werkwijze behelst het toepassen van de probabilistische segmentatiemethode op de hyperspectrale beelden van Schiermonnikoog. Hiervoor zijn de C-programma's en bash-scripts die door B. Gorte geschreven zijn gebruikt. In veel gevallen moesten deze aangepast worden voor ze toepasbaar waren voor hyperspectrale data of voor natuurlijke vegetatie. Verder zijn aanvullende routines geschreven en zijn de losse routines aan elkaar gekoppeld zodat de gehele procedure met een enkel commando uitgevoerd kan worden.

Deze werkzaamheden zijn aangevuld met een bestudering van de beschikbare data en een beknopt literatuuronderzoek. Tijdens het literatuuronderzoek ontstond het idee voor een nieuwe methode voor fuzzy training voor maximum-likelihood-classificatie. Deze methode is geprogrammeerd. Hoewel een test aangaf dat de methode zou kunnen werken, is deze uiteindelijk niet in het classificatieprogramma opgenomen. De reden daarvoor is dat uiteindelijk niet maximum-likelihood-classificatie maar k-nearest-neighbours-classificatie gebruikt is. Bovendien bleek de gebruikte data van Schiermonnikoog zich niet zo voor fuzzy training te lenen.

Het operationaliseren en aanpassen van software vergt veel tijd. Daardoor ligt de nadruk van het afstudeeronderzoek op het uitbreiden van de probabilistische segmentatiemethode voor toepassing op de HyMap-beelden van Schiermonnikoog en het programmeren hiervan. Het onderzoek is afgesloten met een aantal experimenten. Hierbij is voor verschillende klassen-indelingen de aangepaste methode van probabilistische segmentatie uitgevoerd en vergeleken met classificatie zonder probabilistische segmentatie.

1.4. Leeswijzer

In hoofdstuk 2 zal eerst achtergrondinformatie gegeven worden over mixed pixels en fuzzy classificatie. Vervolgens zullen in hoofdstuk 3 de HyMap-beelden en de daarbij behorende veldgegevens waar de methode op getest is beschreven worden. In hoofdstuk 4 wordt besproken hoe in het geval van fuzzy veldgegevens de training voor maximum-likelihood-classificatie er uit zou moeten zien. De probabilistische segmentatie wordt in hoofdstuk 5 behandeld, alsmede de gemaakte aanpassing voor natuurlijke vegetatie. De keuze voor een classificatiemethode voor het schatten van de noodzakelijke kansen wordt in hoofdstuk 6 gemaakt. De resultaten van de toepassing van probabilistische segmentatie worden geëvalueerd in hoofdstuk 7. Conclusies en aanbevelingen naar aanleiding van het onderzoek zijn tenslotte te vinden in hoofdstuk 8.

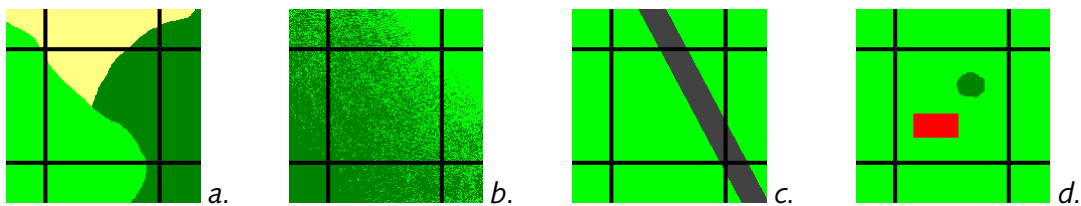
De resultaten van de toepassing van probabilistische segmentatie op de beelden van Schiermonnikoog zijn te vinden in bijlage 1. In bijlage 2 is een handleiding voor het gebruikte programma voor probabilistische segmentatie opgenomen.

2. Mixed pixels en fuzzy classificatie

2.1. Mixed pixels

Een digitaal beeld is opgebouwd uit pixels, welke voorgesteld worden als vierkante beeldvlakjes, die in een regelmatig grid het volledige oppervlak bedekken. De spectrale informatie van een pixel is afkomstig uit een gebied dat bedekt wordt door de "instantaneous field of view" (IFOV) van de sensor. De traditionele methoden voor het classificeren van remote-sensing-beelden kennen op grond van deze spectrale informatie aan ieder pixel één klasse toe. In veel gevallen is het echter zo dat in de werkelijkheid die het pixel beschrijft niet slechts één klasse voor komt. Een pixel waarin in meer of mindere mate twee of meer klassen voorkomen, wordt een mixed pixel genoemd. Er kunnen vier typen mixed pixels onderscheiden worden op grond van de oorzaak van de aanwezigheid van meer dan één klasse [Fisher 1997 blz. 681]:

- Grens tussen twee of meer klassen (bijv. een bosrand);
- Geleidelijke overgang tussen klassen (bijv. "ecotone");
- Lineaire subpixel-objecten (bijv. een weg);
- Kleine subpixel-objecten (bijv. een huis of boom).



Figuur 2.1. De vier typen van mixed pixels [Fisher 1997 blz. 683].

Het is niet meteen helemaal duidelijk wat het overgangspixel (type b) nu eigenlijk inhoudt. Fisher [1997 blz. 681] zegt dat hier sprake is van een overgang (intergrade) tussen centrale concepten van karteerbare fenomenen. Het lijkt er dus om te gaan dat in de andere drie gevallen de grens tussen de klassen in de werkelijkheid wel duidelijk is, maar dat deze bij type b in het terrein ook onduidelijk is. Dat hangt echter nogal af van de schaal waarop je kijkt. Als je maar genoeg inzoomt wordt de scherpe grens van een bos, weg of zelfs huis vaag. Een punt of grens heeft een geometrische precisie waarmee je hem in het terrein kunt aanwijzen. Dit wordt in de landmeetkunde de idealisatieprecisie genoemd.

Wat met een bepaalde ruimtelijke resolutie nog een bepaald type mixed pixel is, kan bij een hogere of lagere resolutie wel eens een ander type zijn. Een voorbeeld hiervan is een grens tussen bos en hei, die bij pixels van 1000 meter grenspixels (type a) heeft. Bij pixels van 100 meter kan de bosrand wel eens niet zo scherp meer blijken te zijn, door een langzaam afnemende dichtheid van de bomen. In de overgangszone komen dan pixels van type b voor. Bij pixels van 10 meter kan één afzonderlijke boom een subpixel-object zijn (type d). Bij pixels van 1 meter kan de grens tussen een boom en de hei weer een pixel van type a opleveren.

Bovenstaand voorbeeld illustreert ook dat mixed pixels niet los staan van klassendefinities. Want bij welke dichtheid van bomen is er nog sprake van bos? Wat je in het trainingsstadium van een classificatie benoemt als puur bos, bepaalt wat pure en wat mixed pixels zijn.

Een remote-sensing-beeld zal altijd een combinatie van pure en mixed pixels bevatten. Het aantal mixed pixels is afhankelijk van de ruimtelijke complexiteit van de werkelijkheid, de ruimtelijke resolutie en de gebruikte klassendefinities.

2.2. Fuzzy classificatie

Mixed pixels zijn een complicatie bij classificatie. Classificatiemethoden die alle pixels op grond van een beslisregel in één van de klassen indelen, worden crisp of harde classificatiemethoden genoemd. Deze definiëren eenduidig wat nog tot welke klasse behoort en wat niet meer. Soms is het echter beter om de vaagheid niet weg te werken, maar deze te behouden (zie paragraaf 1.1) door het pixel in meer of mindere mate in te delen in meerdere klassen. Methoden die dit doen worden fuzzy classificatiemethoden genoemd.

Het woord fuzzy betekent zoveel als vaag. De oorsprong van de term ligt bij de introductie van fuzzy sets als generalisatie van de Booleaanse logica. Hierbij hoeft een bewering niet langer juist of onjuist te zijn, oftewel 1 of 0, maar kan deze ook alle waarden daartussen aannemen. Naast dat fuzzy vaag betekent, is het in het geval van fuzzy classificatie ook een beetje een vage term: er wordt niet altijd hetzelfde onder verstaan. In deze scriptie zal de volgende definitie gebruikt worden:

Een fuzzy classificatie is een classificatie waarbij een object niet in slechts één klasse ingedeeld kan worden, maar in meer of mindere mate tot meerdere klassen kan behoren, uitgedrukt in een waarde per klasse.

Dit is een vrij ruime definitie voor fuzzy classificatie. Het komt er op neer dat alle methoden die niet crisp zijn, fuzzy zijn.

Een aan fuzzy classificatie verwante term is subpixelclassificatie. Ook bij subpixelclassificatie wordt een pixel ingedeeld in meerdere klassen. Hierbij gaat het echter om het schatten van aandelen van de verschillende klassen. Terwijl bij fuzzy classificatie de waarden per klasse niet altijd als aandelen geïnterpreteerd mogen worden. Het kan bijvoorbeeld ook om de kans op een klasse gaan of de mate van lidmaatschap. Fuzzy classificatie is volgens bovenstaande definitie dus een ruimer begrip dan subpixelclassificatie. Dit neemt echter niet weg dat er gepoogd kan worden om uit een fuzzy classificatie met kansen per klasse aandelen af te leiden. Er zijn namelijk drie methoden voor subpixelclassificatie [Eastman en Laney 2002]:

- Aandelen schatten uitgaande van lineaire vermenging van de spectra van pure klassen;
- Het afzonderlijk trainen voor verschillende mengverhoudingen;
- Aandelen afleiden uit een fuzzy classificatie.

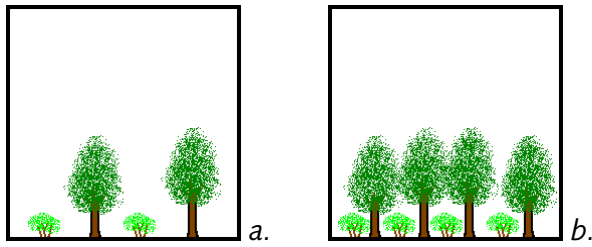
2.3. Lineaire vermenging van spectra

Sommige methoden voor subpixelclassificatie gaan uit van lineaire vermenging van de spectra. Het spectrum van een mixed pixel met half klasse 1 en half klasse 2 ligt dan in de featurespace precies tussen de spectra van de klassen 1 en 2 in. Voor pixels waarin twee klassen naast of door elkaar voorkomen zonder elkaars reflectie te beïnvloeden is deze aanname terecht. Er zijn echter ook mixed pixels met niet-lineair vermengde spectra. Voor mineralen moet voor de verklaring hiervan (op molecuulniveau) gekeken worden naar de interactie van de materialen met opvallend licht [Clark 1999 paragraaf 6.1]. Op het eerste gezicht zou gesteld kunnen worden dat in het geval van vegetatie dit soort zaken zich niet voor doen. Onderstaand voorbeeld toont echter aan dat niet zonder meer waar is. Dit betreft een duidelijke situatie, maar dergelijke effecten kunnen zich ook minder opvallend voordoen.

Voorbeeld

De verhouding van bomen en struiken is in figuur 2.2a hetzelfde als in figuur 2.2b. Toch zal het spectrum van figuur 2.2b meer op dat van bomen lijken dan het spectrum van figuur 2.2a. Als vervolgens een schatting gemaakt zou worden van de

aandelen bomen en struiken zou deze voor figuur 2.2a een ander resultaat geven dan voor figuur 2.2b.



Figuur 2.2a. Een situatie die leidt tot een lineaire vermenging van de spectra van de pure klassen; b. Een situatie die leidt tot een niet-lineaire vermenging van de spectra van de pure klassen.

Met remote sensing is dit probleem niet op te lossen. De gebruiker zal indien hij remote sensing wil gebruiken voor de kartering moeten accepteren dat dergelijke effecten zich voor doen en dat dus figuur 2.2a per definitie een andere verhouding geeft dan 2.2b. De gebruiker moet enige voorzichtigheid betrachten bij het interpreteren van de verhoudingen die uit de remote-sensing-data geschat worden. Deze komen namelijk niet altijd overeen met de manier waarop hij gewend is over verhoudingen te denken.

Een volgend punt van aandacht voor de aanname van lineaire vermenging van spectra is het effect van de sensor. De IFOV (instantaneous field of view) is in werkelijkheid meestal cirkelvormig, in tegenstelling tot het pixel dat als vierkant weergegeven wordt. Daarnaast nemen sommige sensoren meer licht op in het midden van de IFOV dan aan de rand. Bij andere sensoren hebben de IFOV overlap met die van aangrenzende pixels, of er zijn stukken van het terrein binnen geen enkel IFOV vallen. Het effect van de reflectie van een deel van de IFOV op de geregistreerde waarde voor de pixel is onduidelijk, zeker als deze reflectie sterk afwijkt van die van de rest van de IFOV [Fisher 1997 blz. 680].

Hoewel er dus verschillende oorzaken mogelijk zijn voor niet-lineaire vermenging van spectra, zal er in deze scriptie in hoofdstuk 4 uitgegaan worden van lineair samengestelde mixed pixels. Men moet echter in de gaten houden, dat naast meetruis en het effect van de atmosfeer er sprake kan zijn van systematische effecten als gevolg van niet-lineairiteit.

3. Beschrijving van het studiegebied en de data

3.1. Studiegebied

Schiermonnikoog is een van de Nederlandse Waddeneilanden, ontstaan door de depositie van sediment uit de Rijn. Langs de noord- en westkust bevinden zich duinen. Om de polder, met agrarisch gebied en het dorp Schiermonnikoog, te beschermen, worden de duinen op de westelijke helft van het eiland kunstmatig versterkt en ligt er langs de zuidkust een zeedijk. Op het oostelijk deel van het eiland kan de zee door openingen in de duinenrij en de afwezigheid van een zeedijk het land binnen dringen. Het daar gelegen natuurgebied, dat voor het grootste gedeelte uit kwelder bestaat, is het onderwerp van dit onderzoek. Er is voor dit gebied gekozen omdat hiervan hyperspectrale beelden beschikbaar waren in combinatie met veldgegevens.

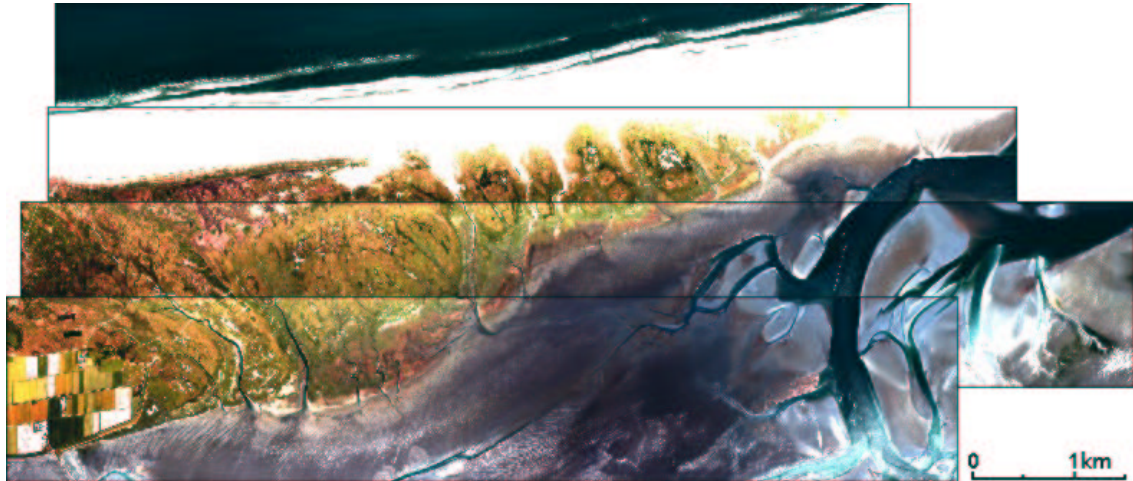


Figuur 3.1. Locatie van het studiegebied

3.2. Beeldmateriaal

Het gebruikte beeldmateriaal (zie figuur 3.2) is ingewonnen met het HyMap-systeem (Hyperspectral Mapping System). Dit is een hyperspectrale whiskbroom-scanner (= across-track-scanner) voor in een vliegtuig. De strookbreedte wordt bepaald door de openingshoek van de scanner (60°), hierbinnen worden 512 pixels opgenomen. De scanner onderscheidt 128 banden van 10 - 20 nm breed in het golflengte-interval van 400 nm tot 2500 nm met een radiometrische resolutie van 16 bits (65536 grijswaarden). De ruimtelijke resolutie van de scanner is afhankelijk van de vlieghoogte en -snelheid alsmede de integratietijd van de sensor. De vlucht vond plaats op 27 mei in 1999. Er zijn vier stroken gevlogen in oost-westrichting. De vlieghoogte van ca. 1650 meter resulteerde in een strookbreedte van een kleine twee kilometer en een pixelgrootte van zo'n 3,5 bij 3,5 meter. Het DLR (Deutsches Zentrum für Luft- und Raumfahrt) heeft de absolute en relatieve oriëntering van de stroken en de geometrische correcties (geo-referentie) bepaald met behulp van differentiële GPS-metingen en oriëntatiewaarnemingen in het vliegtuig. De resulterende geometrische precisie van de beelden is niet gekwantificeerd, maar de standaardafwijking is zeker groter dan een pixel [Skidmore e.a. 2001 blz. 13-14]. Dit maakt het toepassen van een fuzzy trainingsmethode (zie hoofdstuk 4) op deze beelden onmogelijk. Dit werd echter pas tijdens het schrijven van hoofdstuk 4 duidelijk.

Deze HyMap-beelden zijn eerder gebruikt voor het expertsysteem dat bij het ITC in samenwerking met de MD ontwikkeld is [Skidmore e.a. 2001]. Tijdens het onderzoek dat in deze scriptie beschreven wordt, zijn echter de niet geometrisch getransformeerde beelden gebruikt. Wel is er gebruik gemaakt van de veldgegevens die voor dit expertsysteem verzameld zijn.



Figuur 3.2. Compositie van de vier HyMap-stroken. De banden 5, 9, en 19 gebruikt, resulterend in een redelijk natuurgetrouwe afbeelding van de kleuren.

Om eenvoudig gebruik te kunnen maken van trainingssamples in alle stroken zijn de vier stroken achter elkaar geplakt, resulterend in één beeld van 10804 bij 512 pixels (zie figuur 3.3). Omdat de software alleen banden met één byte per pixel aan kan, zijn banden van twee bytes per pixel teruggebracht naar één byte per pixel. Hiervoor is gebruik gemaakt van een stretchmethode. Om niet te veel radiometrische resolutie door uitschieters te verliezen, zijn de pixels met de laagste en de hoogste reflecties respectievelijk op de waarden 0 en 255 gesteld. Dit is gedaan voor de pixels die in het cumulatieve histogram onder de 1% of boven de 99% vallen. Deze stretch is per band afzonderlijk uitgevoerd, maar voor alle vier de stroken tegelijk. Hierdoor verandert de spectrale signatuur van de klassen. Zolang de trainingssamples uit het te classificeren beeld gehaald worden, heeft dit echter geen invloed op de classificatie. Het voordeel van een stretch per band is dat het spectrale onderscheid tussen de klassen zo goed mogelijk bewaard blijft bij het terugbrengen naar één byte per pixel.



Figuur 3.3. De aan elkaar geplakte stroken 4, 3, 2, en 1.

3.3. Veldgegevens

Voor de training en validatie van de classificatie is gebruik gemaakt van de veldgegevens die in mei en juni van het jaar van de beeldopname ingewonnen zijn. In totaal werden op 384 plaatsen in het veld vegetatiegegevens verzameld (veldopnameplaatsen). Deze opnamen bestonden uit gebiedjes van vijf bij vijf meter waar de vegetatiesamenstelling is bepaald. De posities van de opnamen zijn ingemeten met behulp van differentiële GPS. De opnameplaatsen zijn min of meer random verdeeld over het natuurgebied. Er is in het terrein gezocht naar gebiedjes die representatief zijn voor een bepaalde vegetatieklasse, en het liefst zo dat de directie omgeving van het opnamegebied ook tot die klasse behoort.

Er zijn twee typen veldopnamen ingewonnen: gewone veldopnamen en snelle veldopnamen. Voor de 208 gewone veldopnamen zijn gedetailleerde soortensamenstelling en oppervlakte-

schattingen van de aanwezige soorten bepaald. Bij 122 van deze gewone veldopnamen zijn veldspectrometermetingen uitgevoerd. Voor de 176 snelle veldopnamen zijn alleen de samenstelling en de oppervlakteschattingen bepaald van de dominerende soorten en van de soorten die van belang zijn voor het vaststellen van het vegetatietype (indicator species). Bovenstaande informatie is ontleend aan Skidmore e.a. [2001 blz. 12].

De veldopnamen moeten gebruikt worden voor zowel het trainen van de classificatie als het valideren van het resultaat. Hiervoor zijn de veldopnamen gesplitst in een trainingsset en een validatieset (ook wel testset). Als trainingssamples zijn de gewone veldopnamen gebruikt en voor de validatie de snelle veldopnamen. Deze worden als voldoende onafhankelijk beschouwd.

Voor de klassen strand en water zijn nauwelijks of geen veldopnamen gemaakt, omdat deze duidelijk in het beeld herkenbaar zijn. De samples voor deze klassen zijn dan ook in het beeld opgegeven. Voor beide klassen is een groep van 25 aaneengesloten pixels aangewezen voor de trainingsset. Op een andere locatie en in een andere strook zijn voor beide klassen een groep van eveneens 25 pixels aangewezen voor de validatie.

Om de locaties van de veldopnameplaatsen in het beeld te vinden is de inverse geometrische transformatie toegepast op de GPS-coördinaten van de veldopnamen, zodat beeldcoördinaten voor de ongetransformeerde beelden verkregen werden. Vermoedelijk als gevolg van de grote standaardafwijking voor de geometrische precisie van het beeld kwamen de beeldspectra van een aantal veldopnamen niet overeen met de veldgegevens. Daarom zijn alle veldopnamen handmatig gecontroleerd en indien mogelijk gecorrigeerd door een aangrenzend pixel als het juiste te identificeren. Dit tijdrovende werk was gelukkig al uitgevoerd ten behoeve van het expertsysteem.

Punten die in de overlap van twee stroken vallen zijn slechts in één strook teruggezocht. Tijdens de classificatie met de methode die in deze rapportage beschreven wordt, bleek voor een aantal van deze punten niet het juiste strooknummer geregistreerd te zijn. Dit openbaarde zich doordat het wad als vegetatie geclassificeerd werd. Met behulp van de GPS-coördinaten konden deze punten eenvoudig geïdentificeerd worden en kon de juiste strook bepaald worden.



Figuur 3.4. Inwinnen van de veldgegevens

3.4. Klassenindeling

De veldopnamen zijn ingedeeld in 40 subtypen. Deze subtypen kunnen vervolgens in klassen samengevoegd worden. Tijdens het onderzoek zijn vier verschillende klassenindelingen gebruikt. Drie daarvan zijn door landschapsecoloog H. Kloosterman gekozen op grond van verschillende mate van aggregatie (zie onderstaande tabel).

Klassenindelingen met verschillende mate van aggregatie			
Niveau 1	Niveau 2	Niveau 3	Subtypen
Kale grond en ijle vegetatie	Kale grond en ijle vegetatie	Kale grond en ijle vegetatie	O + 1.1 + 2.1
Pionierkwelder	Salicornia-Suaeda-vegetatie	Salicornia (high cover)	2.2
		Salicornia-Suaeda	3.1
		Suaeda-Limonium	3.2
Lage kwelder	Puccinellia-vegetatie	Puccinellia-Salicornia-Limonium	4.1 + 4.2
		Puccinellia-Glaux	4.3 + 4.4
	Limonium-vegetatie	Limonium	5.1
	Halimione-vegetatie	Halimione	6.1 + 6.2
Middelhoge kwelder	Juncus-ger.-vegetatie	Juncus-Suaeda-Limonium	8.1 + 8.2
		Juncus-Festuca	8.3 + 8.4
	Festuca-vegetatie	Festuca	9.1 + 9.2
		Artemisia	9.3abc
	Juncus-mar.-vegetatie	Juncus mar.	10.1 + 10.2
Hoge kwelder	Elymus-vegetatie	Elymus	11.1 + 11.2
	Potentilla-vegetatie	Potentilla-Juncus	12.1
		Agrostis-Potentilla	12.2 + 12.3 + 12.4
	Sagina-vegetatie	Sagina-Armeria	13.1
Overgang duin kwelder	Overgang duin kwelderveg.	Overgang duin kwelder	14.1 + 15.1 + 15.2
Duin	Duin	Duin	D
Strand	Strand	Strand	B
Water	Water	Water	W

De vierde indeling die gebruikt is heeft dezelfde klassen als voor het expertsysteem gebruikt zijn (zie onderstaande tabel).

Klassenindeling zoals gebruikt in Skidmore e.a. [2001]	
Naam	Afkorting
Water	Water
Kale grond en ijle vegetatie	Bare(SpaSal)
Suaeda-Salicornia-vegetatie	SuSa
Salicornia-Limonium-vegetatie	SalLim
Limonium-vegetatie	Lim
Puccinellia-vegetatie	Pucc
Halimione-vegetatie	Hal
Festuca-Artemisia-vegetatie	FestArte
Juncus-Suaeda-Limonium-vegetatie	JunSuaLim
Juncus-Artriplex-Festuca-vegetatie	JunAtrFes
Festuca-vegetatie	Fest
Juncus mar.-vegetatie	JunM
Elymus-vegetatie	Ely
Agrostis-vegetatie	Agros
Agrostis-Lotus-vegetatie	AgrosLot
Sagina-Armeria-vegetatie	SagArm
Duinvoet	Dunefoot
Duin	Dune
Strand	Beach

Het indelen van de veldopnamen in subtypen en het samenvoegen van deze subtypen in klassen zijn classificaties op zichzelf. Het opstellen van deze indelingen is handmatig geschied, hierdoor zijn deze redelijk subjectief. De subtypen zijn op zichzelf namelijk ook weer mengsels van verschillende plantensoorten, waarbij één soort in meerdere subtypen voor kan komen. Dit resulteert in klassen met een relatief grote spectrale overlap.

4. Fuzzy training voor maximum-likelihood-classificatie

Training is de eerste stap van een supervised classificatie. Het doel van de training is op te geven welke klassen er onderscheiden dienen te worden, en wat de kenmerken van deze klassen zijn. Om aan de spectrale kenmerken van de klassen te komen wordt bij vegetatie-classificatie meestal geen gebruik gemaakt van bekende spectra uit een bibliotheek. De atmosferische omstandigheden en het groeistadium van de vegetatie ten tijde van de opname beïnvloeden de spectra namelijk sterk. Daarom worden voor iedere klasse een aantal representatieve pixels in het beeld als trainingssamples geselecteerd.

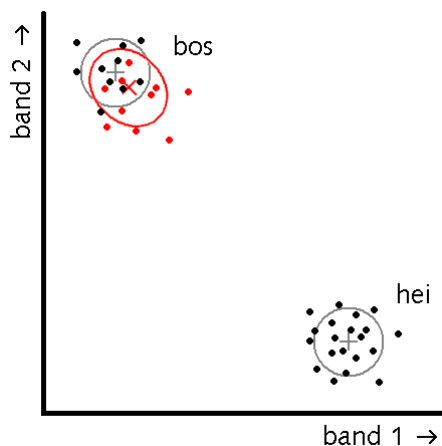
4.1. Aanleiding

De meeste bestaande technieken hebben pure pixels nodig voor de training. Dit maakt de training voor fuzzy classificatie van natuurlijke vegetatie lastig. Want de reden om voor fuzzy classificatie te kiezen is juist dat men uit ervaring weet dat er bij natuurlijke vegetatie veel mixed pixels in het terrein voorkomen. Hierdoor wordt het moeilijk om voldoende pure trainingssamples te vinden.

Een van de mogelijkheden om dit probleem het hoofd te bieden is bijna pure pixels te gebruiken voor de training. Er zijn zelfs methoden om de puurste pixels in een beeld te vinden, door te zoeken naar uitersten in de featurespace zoals de Pixel Purity Index (PPI) [ENVI Tutorials 1999 blz. 267]. Het gebruik van bijna pure trainingssamples is echter suboptimaal. Hierbij wordt namelijk een deel van de geschatte spreiding van een klasse niet veroorzaakt door de natuurlijke variatie in het spectrum van die klasse, maar door de mate waarin er aandelen van andere klassen aanwezig zijn. Bij crisp classificatie zal dit niet zo snel tot problemen leiden, omdat voor de meeste pixels de meest waarschijnlijke klasse nog steeds de goede klasse zal zijn. Bij fuzzy classificatie wil men echter de kansen (of een andere fuzzy waarde) als eindresultaat geven en die zullen op deze manier een systematische fout bevatten.

Voorbeeld

Stel, we willen de klassen bos en hei onderscheiden. We veronderstellen dat deze beide een gelijke spreiding hebben en dat de banden onderling ongecorreleerd zijn. (zie figuur 4.1).



Figuur 4.1. Featurespace met pure samples van de klasse hei en slechts deels pure samples van de klasse bos.

Voor de training van de klasse hei beschikken we over 20 pure samples, maar voor de klasse bos gebruiken we 10 pure samples en 10 samples waar 10% hei in voor komt. In dat geval zal de geschatte spreiding (de ellips in figuur 4.1) van de klasse bos groter zijn dan de werkelijke (de cirkel). Hierdoor zal een pixel met half bos en half hei geclassificeerd worden als zijnde meer bos dan hei.

De oplossing voor een gebrek aan pure trainingssamples moet niet gezocht worden in het gebruik van bijna pure samples alsof deze puur zijn, maar in het schatten van pure spectra uit mixed trainingssamples. In de ideale situatie zijn de training, de classificatie en de validatie fuzzy [Foody 1999].

4.2. Nieuwe methode voor fuzzy training

Bij maximum-likelihood-classificatie wordt voor iedere klasse uit de trainingsset de gemiddelde grijswaarde per spectrale band berekend en de spreiding en correlatie met de andere banden. Hierbij wordt de normale verdeling verondersteld. Hoe kunnen deze kenmerken geschat worden met mixed pixels? F. Wang [1990] stelt een methode voor waarbij een gewogen gemiddelde en een gewogen empirische (co)variantie worden berekend, met de fractie van de betreffende klasse als gewicht. Deze procedure leidt echter niet tot zuivere schattingen voor de spectra van de pure klassen [Eastman en Laney 2002 blz. 1151]. De niet-pure trainingssamples trekken namelijk (ondanks hun lagere gewicht) het gewogen gemiddelde in de richting van de niet-pure spectra.

Een andere manier om mixed pixels als trainingssamples te gebruiken gaat uit van het afzonderlijk trainen voor verschillende mengverhoudingen. Hiervoor zijn echter trainingssamples voor alle te onderscheiden verhoudingen benodigd. Hiervoor zouden erg veel trainingssamples nodig zijn wat zeer omvangrijk veldwerk met zich mee zou brengen.

In deze paragraaf wordt daarom een nieuwe methode beschreven waarmee zuivere schattingen van pure spectra en hun (co)variantie verkregen kunnen worden uit mixed trainingssamples. Hierbij wordt uitgegaan van lineaire vermenging van de spectra (zie paragraaf 3.3). Voor de grijswaarde van een mixed pixel geldt dan:

$$\chi_{b_j s_i} = f_{C_1 s_i} \cdot \chi_{b_j C_1} + f_{C_2 s_i} \cdot \chi_{b_j C_2} + \dots = \sum_{k=1}^K f_{C_k s_i} \cdot \chi_{b_j C_k}$$

met: $\chi_{b_j s_i}$ grijswaarde voor band j van sample i ;
 $f_{C_k s_i}$ fractie van klasse k in sample i ;
 $\chi_{b_j C_k}$ gemiddelde grijswaarde voor band j van klasse k ;
 K het aantal klassen.

De grijswaarde $\chi_{b_j s_i}$ is element j van de featurevector van sample i . Een featurevector is een vector met de spectrale signatuur van een pixel. Deze wordt hier aangeduid met de χ (chi), de eerste letter van het Griekse woord chroma dat kleur betekent.

De som van alle fracties moet 1 zijn, dus geldt tevens:

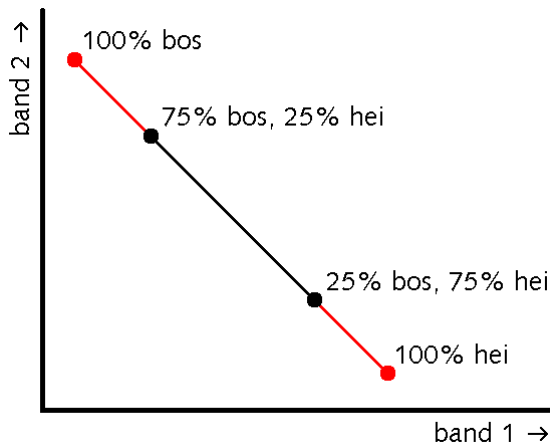
$$\sum_{k=1}^K f_{C_k s_i} = 1$$

Volgens bovenstaand model liggen in de featurespace punten met verschillende mengverhoudingen van twee klassen op een rechte lijn tussen de punten van de twee pure spectra

in. De positie op deze lijn is recht evenredig met de mengverhouding. Dit maakt het mogelijk om uit gemengde trainingssamples de spectra van de pure klassen te bepalen, mits de mengverhoudingen van de samples bekend zijn (zie onderstaand voorbeeld).

Voorbeeld

Stel dat we twee samples hebben, één met 25% bos en 75% hei en één met 75% bos en 25% hei. Als we in de featurespace een lijn trekken tussen deze twee samples en deze aan beide zijden nog eens de helft van z'n lengte doortrekken, dan komen we uit op schattingen voor 100% samples bos en hei (zie figuur 4.2).



Figuur 4.2. Featurespace met schatting van twee pure spectra uit twee mixed samples.

Als we echter niet over twee maar over vele samples beschikken, zullen deze nooit exact op een lijn liggen, zodat zij niet meer eenduidig een waarde voor de pure spectra vastleggen. Met behulp van de kleinstekwadratenvereffening kan dan een schatting van de pure spectra verkregen worden (zie figuur 4.3).

4.2.1. Kleinstekwadratenschatting van pure spectra

Voor het schatten van de spectra van de pure klassen kan gebruik gemaakt worden van de vereffeningstheorie zoals beschreven in Teunissen [1994]. Bovenstaande formules geschreven als waarnemingsvergelijkingen $E\{y\} = A(x)$ zijn:

$$\begin{aligned}
 E\{\chi_{b_j s_i}\} &= \sum_{k=1}^{K-1} (f_{C_k s_i} \cdot \chi_{b_j C_k}) + \left(1 - \sum_{k=1}^{K-1} (f_{C_k s_i})\right) \cdot \chi_{b_j C_K} \\
 E\{f_{C_1 s_i}\} &= f_{C_1 s_i} \\
 &\vdots \\
 E\{f_{C_{K-1} s_i}\} &= f_{C_{K-1} s_i}
 \end{aligned}$$

Hierbij zijn de waarnemingen y de spectra van de trainingssamples en de klassenaandelen van deze samples. De onbekenden x zijn de spectra van de pure klassen en de klassenaandelen van de trainingssamples. Deze vergelijkingen zijn echter niet lineair en zullen voor de kleinstekwadratenvereffening gelineariseerd moeten worden. De gelineariseerde vergelijkingen zijn:

$$\begin{aligned}
E\{\Delta \chi_{b_j s_i}\} &= \sum_{k=1}^{K-1} \left(f_{C_k s_i}^o \cdot \Delta \chi_{b_j C_k} \right) + \left(1 - \sum_{k=1}^{K-1} f_{C_k s_i}^o \right) \cdot \Delta \chi_{b_j C_K} + \sum_{k=1}^{K-1} \left((\chi_{b_j C_k}^o - \chi_{b_j C_K}^o) \cdot \Delta f_{C_k s_i} \right) \\
E\{\Delta f_{C_1 s_i}\} &= \Delta f_{C_1 s_i} \\
&\vdots \\
E\{\Delta f_{C_{K-1} s_i}\} &= \Delta f_{C_{K-1} s_i}
\end{aligned}$$

met: de waarnemingen $\Delta y = y - y^o$ en de onbekenden $\Delta x = x - x^o$, waarbij x^o benaderde waarden voor de onbekenden x zijn (bijvoorbeeld bepaald uit de puurste samples) en y^o verkregen wordt uit $y^o = A(x^o)$.

De gelineariseerde vergelijkingen zijn in matrixvorm te herschrijven als $E\{\Delta y\} = A \cdot \Delta x$. De kleinstekwadratenoplossing van de onbekenden (spectra en klassenaandelen) volgt dan uit:

$$\hat{x} = x^o + (A^T Q_y^{-1} A)^{-1} A^T Q_y^{-1} \Delta y \quad \text{met: } Q_y = D\{y\}$$

Bij niet heel nauwkeurige benaderde waarden zal het nodig zijn om $\hat{x}$ iteratief te schatten. Door een aantal maal de uitkomst als nieuwe benaderde waarde te gebruiken, kan een voldoende nauwkeurige schatting verkregen worden.

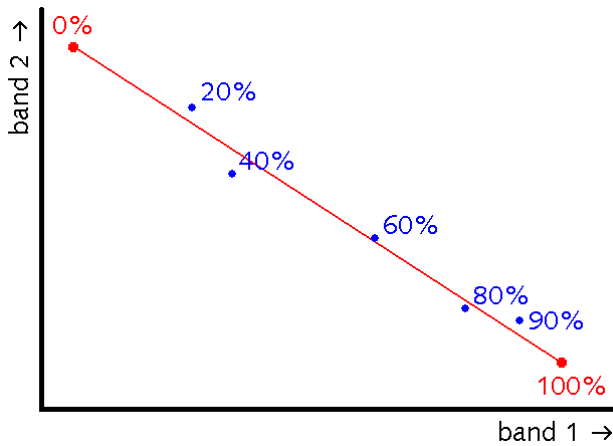
Voorbeeld

Voor twee klassen, twee spectrale banden en vijf samples wordt het uitgeschreven gelineariseerde model van waarnemingsvergelijkingen:

$$E\left\{ \begin{bmatrix} \Delta \chi_{b_1 s_1} \\ \Delta \chi_{b_2 s_1} \\ \Delta f_{C_1 s_1} \\ \vdots \\ \Delta \chi_{b_1 s_5} \\ \Delta \chi_{b_2 s_5} \\ \Delta f_{C_1 s_5} \end{bmatrix} \right\} = A \cdot \begin{bmatrix} \Delta \chi_{b_1 C_1} \\ \Delta \chi_{b_2 C_1} \\ \Delta \chi_{b_1 C_2} \\ \Delta \chi_{b_2 C_2} \\ \Delta f_{C_1 s_1} \\ \vdots \\ \Delta f_{C_1 s_5} \end{bmatrix}$$

met:

$$A = \left[\begin{array}{cccc|cccc} f_{C_1 s_1}^o & 0 & 1 - f_{C_1 s_1}^o & 0 & \chi_{b_1 C_1}^o - \chi_{b_1 C_2}^o & & & 0 \\ 0 & f_{C_1 s_1}^o & 0 & 1 - f_{C_1 s_1}^o & \chi_{b_2 C_1}^o - \chi_{b_2 C_2}^o & & & 0 \\ 0 & 0 & 0 & 0 & 1 & & & 0 \\ \vdots & \vdots & \vdots & \vdots & & \ddots & & \\ f_{C_1 s_5}^o & 0 & 1 - f_{C_1 s_5}^o & 0 & 0 & & \chi_{b_1 C_1}^o - \chi_{b_1 C_2}^o & \\ 0 & f_{C_1 s_5}^o & 0 & 1 - f_{C_1 s_5}^o & 0 & & \chi_{b_2 C_1}^o - \chi_{b_2 C_2}^o & \\ 0 & 0 & 0 & 0 & 0 & & 1 & \end{array} \right]$$



Figuur 4.3. Featurespace met kleinstekwadratenschatting van twee pure spectra uit vijf mixed pixels. De percentages geven het aandeel van één klasse weer, het aandeel van de andere klasse is 100% min dit percentage.

Er is verondersteld dat de klassenaandelen zo ingewonnen zijn dat ze voor ieder trainingssample samen 1 geven. Daarom is het aandeel van de klasse K voor ieder sample geen onafhankelijke waarneming en is deze niet opgenomen in de waarnemingsvergelijkingen. Wanneer de aandelen niet altijd samen 1 geven, is het aandeel van klasse K wel een onafhankelijke waarneming, maar nog steeds geen onafhankelijke onbekende. In dat geval komt er voor ieder trainingssample s_i één waarnemingsvergelijking bij:

$$E\{f_{C_K s_i}\} = 1 - \sum_{k=1}^{K-1} f_{C_k s_i}$$

Het onderkennen van de aanwezigheid van een onbekende klasse is in principe mogelijk. Deze zou dan ook geschat moeten worden. Zeer waarschijnlijk is echter de veronderstelling van een normale verdeling voor deze onbekende klasse onterecht.

De covariantiematrix van de schatting van de onbekenden $\hat{x}$ kan berekend worden met:

$$Q_{\hat{x}} = (A^T Q_y^{-1} A)^{-1}$$

Dit is wat anders dan de spreiding van de pure spectra die in de volgende subparagraaf geschat wordt. In het geval van pure pixels laat dit verschil zich eenvoudiger uitleggen. De spreiding van de pure spectra die in de volgende subparagraaf geschat wordt, is de variatie die voorkomt in het spectrum van pixels van een klasse. Dit is de (co)variantie van de waarnemingen (Q_y). De covariantiematrix van de onbekenden ($Q_{\hat{x}}$) geeft aan met welke precisie de onbekenden (o.a. de verwachtingswaarde van de pure spectra) bepaald kunnen worden uit de waarnemingen.

4.2.2. Variantiecomponentenmodelschatting van pure spectra

Voor een maximum-likelihood-classificatie is ook de spreiding van de spectra nodig voor classificatie. Normaal gesproken wordt een schatting van het kansmodel voor een klasse verkregen met de formule voor de empirische (co)variantie:

$$\sigma_{\chi_{b_j C_k} \chi_{b_h C_k}} = \frac{\sum_{i=1}^{N_k} (\chi_{b_j S_i} - \chi_{b_j C_k})(\chi_{b_h S_i} - \chi_{b_h C_k})}{N_k - 1}$$

$$\text{met: } \chi_{b_j C_k} = \frac{\sum_{i=1}^{N_k} \chi_{b_j S_i}}{N_k}$$

N_{C_k} het aantal samples voor klasse k

Bij het bepalen van pure spectra uit mixed pixels is deze formule niet meer toereikend. In dat geval moet de meer algemene methode van variantiecomponentenmodelschatting gebruikt worden, net zoals de formule voor het gemiddelde vervangen moest worden voor een kleinstekwadratenschatting. Onder bepaalde omstandigheden leiden de algemene formules van de kleinstekwadratenvereffening en de variantiecomponentenmodelschatting tot de vereenvoudigde formules van gemiddelde en empirische (co)variantie.

Bij variantiecomponentenmodelschatting, zoals beschreven in Verhoef [1997 blz. 33-53], wordt de covariantiematrix Q_y geschreven als de gewogen som van P componenten Q_p waarbij de gewichten σ_p^2 geschat worden:

$$Q_y = \sigma_1^2 Q_1 + \sigma_2^2 Q_2 + \dots + \sigma_P^2 Q_P = \sum_{p=1}^P \sigma_p^2 Q_p$$

met: σ_p^2 onbekende (co)variantiefactor p ;
 Q_p cofactormatrix p .

De oplossing voor de (co)variantiefactoren volgt uit:

$$\hat{\sigma} = N^{-1} l$$

met: $\sigma = (\sigma_1^2 \dots \sigma_P^2)^T$

$$N_{pq} = \text{spoor} (Q_y^{-1} P_A^\perp Q_p Q_y^{-1} P_A^\perp Q_q)$$

$$l_p = y^T Q_y^{-1} P_A^\perp Q_p Q_y^{-1} P_A^\perp y$$

$$P_A^\perp = I - A \cdot (A^T Q_y^{-1} A)^{-1} A^T Q_y^{-1}$$

Aangezien Q_y , die bepaald wordt, ook in de formule voor N_{pl} voorkomt, zal men over benaderde waarden voor de (co)variantiefactoren moeten beschikken. Hiervoor kan hetzelfde kansmodel gebruikt worden dat voor de vereffening verondersteld is. De berekening zal dan iteratief moeten plaatsvinden, waarbij voor elke iteratiestap de benaderde waarden worden verbeterd.

Opstellen van het model

In principe kunnen de componenten Q_p op veel manieren gekozen worden. Niet iedere keuze is echter even zinvol. We willen, net als bij de aanwezigheid van pure trainingssamples, de spreiding van de pure spectra schatten. Daarom moet de covariantiematrix van een sample met een willekeurige mengverhouding uitgedrukt worden als functie van de covariantiematrix van de pure spectra.

Aangezien voor gemengde pixels gold:

$$E\{\chi_{b_j s_i}\} = \sum_{k=1}^{K-1} \left(f_{C_k s_i} \cdot \chi_{b_j C_k} \right) + \left(1 - \sum_{k=1}^{K-1} f_{C_k s_i} \right) \cdot \chi_{b_j C_K}$$

volgt uit de voortplantingswetten van varianties hoe de covariantiematrix van de samples met willekeurige mengverhouding afhankelijk zijn van de covariantiematrices van de pure spectra:

$$Q_{\chi_{s_i}} = \sum_{k=1}^{K-1} \left(f_{C_k s_i}^2 \cdot Q_{\chi_{C_k}} \right) + \left(1 - \sum_{k=1}^{K-1} f_{C_k s_i} \right)^2 \cdot Q_{\chi_{C_K}} + \dots$$

De termen die de correlatie tussen de klassen beschrijven zijn weggelaten, omdat deze bij crisp maximum-likelihood-training ook niet geschat worden. Dit zijn termen die de covariantie van de grijswaarden van één band (of twee verschillende banden) tussen twee klassen zijn uitdrukken, zoals $\sigma_{\chi_{b_1 C_1} \chi_{b_1 C_2}}$.

De covariantiematrices van de pure spectra $Q_{\chi_{C_k}}$ kunnen vervolgens uitgedrukt worden in P te kiezen componenten Q_p . Wanneer we evenveel parameters van de spreiding willen schatten als bij gewone maximum-likelihood-training, zal voor iedere variantie en iedere covariantie een component opgesteld moeten worden. Dit resulteert in $P = \frac{1}{2}(J + J^2)$ componenten (met J het aantal banden). Om het model compleet te krijgen moet ook een variantie voor de waarnemingen van de klassenaandelen geschat worden. Hoewel deze waarnemingen gecorreleerd zouden kunnen zijn, zullen hier geen factoren voor geschat worden.

Voorbeeld

Bij twee banden en twee klassen geldt voor de covariantiematrix van de samples:

$$Q_{\chi_{s_i}} = f_{C_1 s_i}^2 \cdot Q_{\chi_{C_1}} + (1 - f_{C_1 s_i})^2 \cdot Q_{\chi_{C_2}} + \dots$$

Wanneer de covariantiematrices van de pure klassen $Q_{\chi_{C_k}}$ gemodelleerd worden als:

$$Q_{\chi_{C_k}} = \sigma_{\chi_{b_1 C_k}}^2 \begin{bmatrix} 1 & 0 \\ 0 & 0 \end{bmatrix} + \sigma_{\chi_{b_2 C_k}}^2 \begin{bmatrix} 0 & 0 \\ 0 & 1 \end{bmatrix} + \sigma_{\chi_{b_1 C_k} \chi_{b_2 C_k}} \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix}$$

Dan wordt het variantiecomponentenmodel:

$$Q_y = \sigma_{\chi_{b_1 C_1}}^2 \begin{bmatrix} \ddots & & & \\ & f_{C_1 s_i}^2 & 0 & 0 \\ & 0 & 0 & 0 \\ & 0 & 0 & 0 \\ & & & \ddots \end{bmatrix} + \sigma_{\chi_{b_1 C_2}}^2 \begin{bmatrix} \ddots & & & \\ & (1 - f_{C_1 s_i})^2 & 0 & 0 \\ & 0 & 0 & 0 \\ & 0 & 0 & 0 \\ & & & \ddots \end{bmatrix} +$$

$$\sigma_{\chi_{b_2 C_1}}^2 \begin{bmatrix} \ddots & & & \\ & 0 & 0 & 0 \\ & 0 & f_{C_1 s_i}^2 & 0 \\ & 0 & 0 & 0 \\ & & & \ddots \end{bmatrix} + \sigma_{\chi_{b_2 C_2}}^2 \begin{bmatrix} \ddots & & & \\ & 0 & 0 & 0 \\ & 0 & (1 - f_{C_1 s_i})^2 & 0 \\ & 0 & 0 & 0 \\ & & & \ddots \end{bmatrix} +$$

$$\sigma_{\chi_{b_1 C_1} \chi_{b_2 C_1}} \begin{bmatrix} \ddots & & & \\ & 0 & f_{C_1 s_i}^2 & 0 \\ & f_{C_1 s_i}^2 & 0 & 0 \\ & 0 & 0 & 0 \\ & & & \ddots \end{bmatrix} +$$

$$\sigma_{\chi_{b_1 C_2} \chi_{b_2 C_2}} \begin{bmatrix} \ddots & & & \\ & 0 & (1-f_{C_1 s_i})^2 & 0 \\ & (1-f_{C_1 s_i})^2 & 0 & 0 \\ & 0 & 0 & 0 \\ & & & \ddots \end{bmatrix} + \sigma_{f_{C_1}}^2 \begin{bmatrix} \ddots & & & \\ & 0 & 0 & 0 \\ & 0 & 0 & 0 \\ & 0 & 0 & 1 \\ & & & \ddots \end{bmatrix}$$

De termen als gevolg van de correlatie tussen de klassen zijn wederom weggelaten. Om het model compleet te krijgen worden ook varianties voor de waarnemingen van de klassenaandelen geschat, eventuele covarianties zijn verwaarloosd.

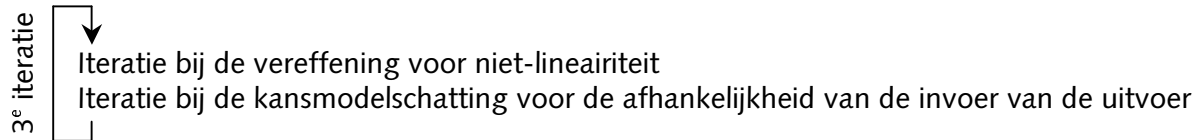
Net als bij crisp training voor maximum-likelihood-classificatie moeten er wel voldoende trainingssamples beschikbaar zijn om alle (co)varianties te schatten. Een negatieve schatting voor een variantiefactor duidt meestal op een slecht gekozen model of te weinig samples. De precisie van de geschatte factoren is namelijk afhankelijk van de overvalligheid. De covariantiematrix van de geschatte (co)varianties kan berekend worden met:

$$Q_{\hat{\sigma}} = 2N^{-1} \quad \text{met: } N_{pq} = \text{spoor} \left(Q_y^{-1} P_A^{\perp} Q_p Q_y^{-1} P_A^{\perp} Q_q \right)$$

Dit geeft de precisie waarmee de precisie geschat is.

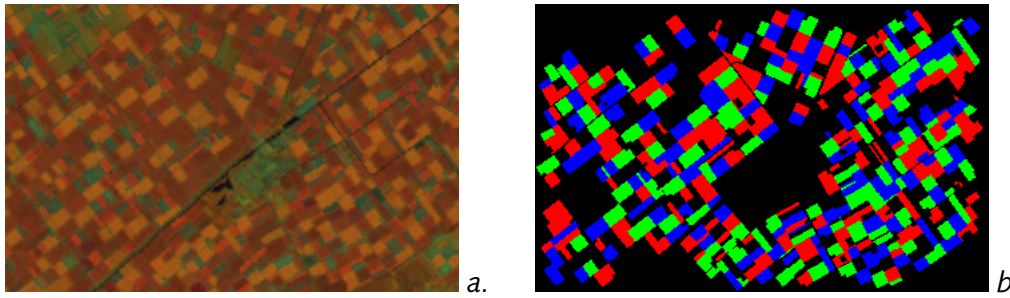
4.2.3. Iteratie

De kansmodelschatting wordt uitgevoerd na de vereffening. Het resultaat van de vereffening is namelijk nodig voor de kansmodelschatting. Omdat het resultaat van de vereffening ook afhankelijk is van het veronderstelde kansmodel zal een iteratie uitgevoerd moeten worden. Dit is een derde iteratie die over de twee eerdergenoemde iteraties uitgevoerd moet worden.



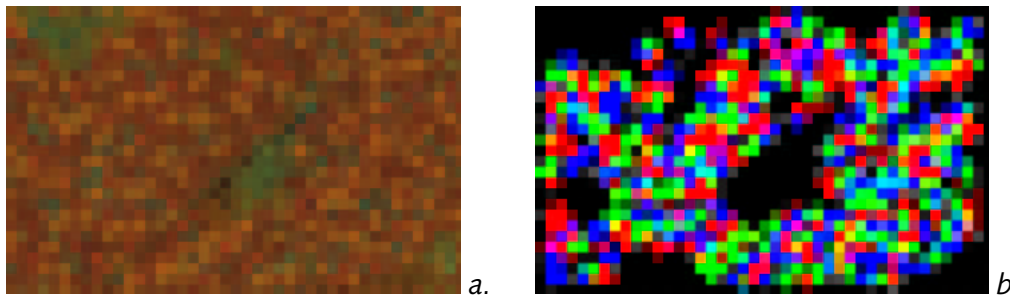
4.3. Experiment

De methode voor fuzzy training is getest door een experiment uit te voeren. Om hierbij aan mixed pixels met een bekende klassenverdeling te komen, zijn een beeld en een bijbehorend rasterbestand met veldgegevens grover geresampled. Het gebruikte beeld is een deel van een opname uit 1987 van de Landsat Thematic Mapper. Dit beeld met pixels van 30 bij 30 meter beslaat een gebied rond Biddinghuizen in Flevoland (zie figuur 4.4a). De banden 4, 5 en 3 zijn gebruikt. Voor dit experiment zijn alleen de drie meest voorkomende gewassen gebruikt: tarwe, aardappels en suikerbieten (zie figuur 4.4b).



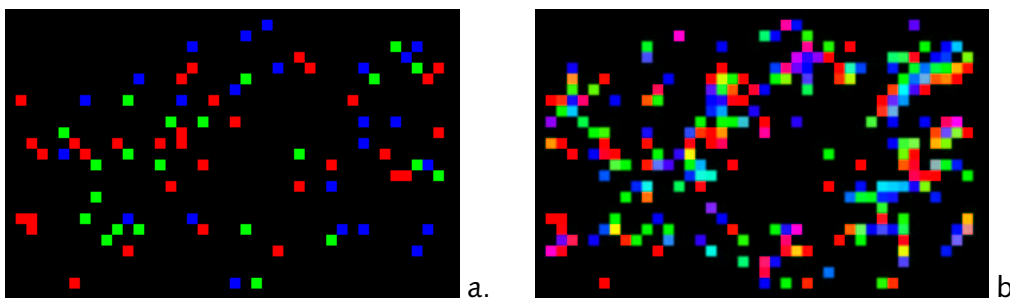
Figuur 4.4a. Landsat-beeld, banden 4 (rood), 5 (groen) en 3 (blauw); b. Rasterbestand met veldgegevens, klassen tarwe (rood), aardappels (groen), suikerbieten (blauw).

Zowel het Landsat-beeld als het rasterbestand met veldgegevens zijn geresampled met een factor acht. In elke band van het beeld is per gebiedje van 8 bij 8 pixels de gemiddelde grijs-waarde genomen, resulterend in drie nieuwe banden met pixels van 240 bij 240 meter (zie figuur 4.5a). In het rasterbestand van de veldgegevens zijn per gebiedje van 64 pixels de aandelen van de drie klassen berekend (zie figuur 4.5b).



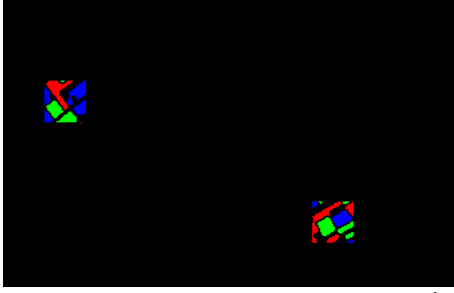
Figuur 4.5a. Geresampled Landsat-beeld alsof het is opgenomen door een sensor met grotere pixels; b. Geresampled rasterbestand met veldgegevens met de oppervlakte-aandelen van de drie gewassen als rood, groen en blauw.

Doordat er in de originele veldgegevens meer klassen aanwezig zijn dan tarwe, aardappels en suikerbieten zijn niet alle geresamplede pixels bruikbaar voor fuzzy training. Voor crisp training zijn echter nog minder samples beschikbaar. Er zijn 243 pixels waarin de drie klassen samen meer dan 90% bedekken en maar 79 pixels waar één klasse 90% bedekt (zie figuur 4.6).



Figuur 4.6a. Voor crisp training bruikbare geresamplede pixels; b. Voor fuzzy training bruikbare geresamplede pixels.

De helft van de veldgegevens (121 van de pixels uit figuur 4.6b) is gebruikt voor fuzzy training op het geresamplede beeld. Deze is vergeleken met crisp training op de pixels uit figuur 4.7 in het ongeresamplede beeld. De verkregen spectra en hun covariantiematrices waren verschillend. De overall-accuracies van classificaties voor fuzzy en crisp geschatte spectra lagen echter beide rond de 80% en verschilden slechts enkele procenten.



Figuur 4.7. Voor crisp training gebruikte ongesampled pixels.

Het grijswaardebereik van het beeld is 6-173. De fuzzy geschatte gemiddelden van de pure klassen verschillen per band ca. 2,4 met de spectra bepaald uit figuur 4.4b. De doorsnee standaardafwijking van deze schattingen $\sigma_{\hat{\chi}_{bjC_k}}$ is 0,79. Het verschil in de geschatte gemiddelde grijswaarden is dus zo'n drie keer de standaardafwijking! De spectra bepaald uit figuur 4.7 verschilden echter ook ca. 2,3 met de spectra bepaald uit figuur 4.4b, en ca. 2,6 met de fuzzy geschatte spectra. Blijkbaar is de afwijking toch niet zo uitzonderlijk groot. De geschatte spreiding (standaardafwijking) in een band voor een klasse $\hat{\sigma}_{\chi_{bjC_k}}$ is zo'n 3,4. De standaardafwijking van deze geschatte standaardafwijking $\sigma_{\hat{\sigma}_{\chi_{bjC_k}}}$ is zo'n 1,05. De spreiding geschat uit figuur 4.4a is zo'n 4,0. Dit komt dus overeen met fuzzy geschatte spreiding.

De standaardafwijking waarmee de klassenaandelen in de vereffening bepaald worden $\sigma_{\hat{f}_{C_k s_i}}$ is ca. 0,065. De schatting van de meetprecisie (standaardafwijking) waarmee de "waarnemingen" van de klassenaandelen gedaan zijn $\hat{\sigma}_{f_{C_k s_i}}$ is met het variantie-componentenmodel geschat op 0,080. De precisie (standaardafwijking) waarmee deze meetprecisie geschat is $\sigma_{\hat{\sigma}_{f_{C_k s_i}}}$ bedraagt 0,0014.

Omdat er geen beperking is opgegeven voor het interval waarbinnen de fracties moeten vallen, worden ook enkele klassenaandelen geschat en die kleiner zijn dan 0 of groter dan 1. De verst buiten het interval gelegen schatting was -0,09. Wanneer het zeker zou zijn dat de mixed pixels volledig samengesteld zijn uit de drie gewassen, is dit niet goed. Bij dit experiment is echter tot 10 procent van een onbekende klasse in een mixed pixel toegestaan. Daarom zijn schattingen tot een tiende kleiner dan 0 of groter dan 1 niet verwonderlijk. Hoe de vereffening uitgevoerd kan worden met een voorwaarden dat de klassenaandelen tussen 0 en 1 liggen, zou onderzocht moeten worden.

Als a priori waarden voor het kansmodel zijn $\sigma_{\chi_{bjC_k}}^o = 4,0$ en $\sigma_{f_{C_k s_i}}^o = 0,10$ genomen. De invloed van dit aangenomen kansmodel blijkt gering te zijn. Bij andere waarden zoals respectievelijk 2,5 en 0,05 verschilden de geschatte spectra minder dan 0,1. Wanneer het benaderde kansmodel echter te veel afwijkt, convergeert de methode een verkeerde kant op. Bij de a priori waarden van respectievelijk 1,7 en 0,16 resulteerde dit in standaardafwijkingen voor de grijswaarden van 10. Bovendien leken de geschatte pure spectra in de verste verte niet meer op die van de bijna pure samples, de afwijkingen liepen in de tientallen.

De voorgestelde nieuwe methode voor fuzzy training blijkt te kunnen werken. Of de methode ook in de praktijk een verbetering kan geven, zou uit verder onderzoek moeten blijken.

4.4. Beschouwing en aanbevelingen

In dit hoofdstuk is beschreven op welke wijze uit mixed pixels de pure spectra geschat kunnen worden inclusief hun spreiding, zodat vervolgens een maximum-likelihood-classificatie uitgevoerd kan worden alsof er wel pure samples beschikbaar waren.

Een voordeel van deze methode is dat deze schattingen geeft zonder systematische fout als gevolg van het gebruik van bijna pure trainingssamples. Een tweede voordeel is dat meer pixels in een beeld bruikbaar zijn als trainingssample. Dit brengt verschillende nieuwe mogelijkheden met zich mee. Eén mogelijkheid is hele gebieden (ondanks heterogeniteit) voor training te gebruiken, waardoor snel meer samples verkregen kunnen worden. Een andere mogelijkheid is trainingssamples op random plaatsen in te winnen, wat de representativiteit wellicht ten goede komt.

Een nadeel is dat voor alle trainingssamples de fracties van de verschillende klassen bekend moeten zijn. Wanneer deze klassenaandelen in het veld ingewonnen worden, betekent dit bovendien dat de trainingssamples met voldoende precisie in het beeld teruggevonden moeten kunnen worden. Een eventueel tweede nadeel is dat er niet langer een eenvoudige richtlijn te geven is voor het aantal samples dat benodigd is voor een goede beschrijving van de klassen. Het benodigde aantal is namelijk niet langer alleen afhankelijk van het aantal klassen en banden maar ook van de klassenverhoudingen van de samples. Met bijna pure samples zullen er minder nodig zijn dan met samples van alle verhoudingen. In het geval dat alle samples dezelfde mengverhouding hebben is het zelfs onmogelijk om pure spectra te schatten. De covariantiematrices van de spectra $Q_{\hat{x}}$ en van de (co)varianties $Q_{\hat{\sigma}}$ zijn een maat voor de kwaliteit van de beschrijving van de pure klassen.

Uiteindelijk is er voor een k-nearest-neighbours-classificatie gekozen in plaats van een maximum-likelihood-classificatie (zie paragraaf 6.2). Daarom is de beschreven methode voor fuzzy maximum-likelihood-training niet verder onderzocht. Voor toepassing van deze methode zijn enige aanbevelingen voor verder onderzoek te geven:

- Er is aangetoond dat de methode zou kunnen werken, maar nog niet of deze in bepaalde situaties ook beter is. Hiervoor is een aanvullend experiment nodig op een beeld waarvoor crisp training niet mogelijk is en waarvoor tijdens veldwerk de klassenaandelen van fuzzy trainingssamples vastgesteld zijn.
- De correlatie tussen de klassen is steeds verwaarloosd. Verder onderzoek zou uit moeten wijzen of deze correlatie aanwezig is en of het verwaarlozen of juist schatten van deze correlatie betere resultaten geeft. Ook andere mogelijke variaties in het opstellen van het variantiecomponentenmodel zouden onderzocht kunnen worden.
- Wellicht is het mogelijk de formules voor de kansmodelschatting te vereenvoudigen voor deze toepassing (net zoals de gebruikelijke formule voor de empirische (co)variantie een vereenvoudiging is). Bij veelvuldig gebruik van deze methode voor fuzzy training zou het de moeite waard kunnen zijn dit uit te zoeken.
- Tenslotte is het toepassen van de toetsingstheorie aan te bevelen. Foute samples kunnen immers grote invloed hebben op de geschatte spectra en (co)varianties. Bovendien kan dan ook het veronderstelde (kans)model getoetst worden, zoals de lineaire vermenging en de normale verdeling van de pure spectra.

5. Probabilistische segmentatie

Probabilistische segmentatie [Gorte 1998] bestaat uit een integratie van beeldsegmentatie en -classificatie. De classificatie vindt plaats op grond van de formule van Bayes (paragraaf 5.1). Met behulp van het resultaat van een statistische classificatiemethode (paragraaf 5.2) worden lokale a priori kansen op de klassen geschat (paragraaf 5.3). Deze lokale a priori kansen kunnen vervolgens gebruikt worden om de classificatie te verbeteren. Hiervoor is een opdeling van het beeld in segmenten nodig, waarbij de segmenten zo veel mogelijk samenvallen met terreinobjecten (paragraaf 5.4). In dit hoofdstuk zal de theorie van probabilistische segmentatie beschreven worden. Een handleiding voor het gebruikte programma is opgenomen in bijlage 2.

5.1. Classificatie

Het doel van een classificatie is het indelen van objecten in klassen. Bij remote sensing is het gebruikelijk om per pixel te classificeren, maar ook groepen pixels (segmenten) zijn mogelijke objecten voor classificatie. Er zijn veel verschillende methoden ontwikkeld met allemaal hun voor- en nadelen. Van belang bij dit onderzoek is dat de methode per pixel, de kans op iedere klasse moet geven. Dit is namelijk noodzakelijk als invoer voor de probabilistische segmentatie. Daarnaast moet de methode instaat zijn om hyperspectrale data te classificeren. Classificatiemethoden die de kans op iedere klasse berekenen om de klasse met de grootste kans te selecteren worden statistische classificatiemethoden genoemd. Er zijn ook classificatiemethoden die op andere gronden classificeren. Wanneer niet de klasse met de hoogste kans, maar de kans per klasse als uitvoer wordt gegeven, wordt een fuzzy statistische classificatie per pixel verkregen. Deze fuzzy classificatie zou als eindresultaat gebruikt kunnen worden. Bij dit onderzoek dient deze echter als invoer voor de volgende stap in de probabilistische segmentatie, namelijk het bepalen van de klassenaandelen per segment (zie paragraaf 5.3).

5.1.1. Formule van Bayes

Een statistische classificatiemethode bepaalt per pixel de meest waarschijnlijke klasse. Hiervoor moet de kans op iedere klasse bekend zijn. Deze kans dat een pixel p met featurevector χ_p tot klasse k behoort, kan berekend worden met de formule van Bayes:

$$P(C_k | \chi_p) = \frac{P(\chi_p | C_k) \cdot P(C_k)}{P(\chi_p)}$$

met: $P(C_k | \chi_p)$ a posteriori kans op een klasse;

$P(\chi_p | C_k)$ conditionele kans op een featurevector;

$P(C_k)$ a priori kans op een klasse;

$P(\chi_p)$ kans op een featurevector.

Voorbeeld

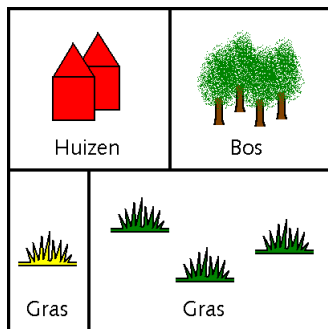
De formule van Bayes laat zich aan de hand van figuur 5.1 eenvoudig uitleggen. Hierin is schematisch een gebied weergegeven met een kwart huizen, een kwart bos en een helft gras. In een remote-sensing-beeld van dit gebied zijn de pixels van de huizen allemaal rood, de pixels van het bos zijn allemaal groen, maar de pixels van het gras zijn voor drie kwart groen en voor een vierde geel. Met de formule van Bayes kan

berekend worden wat de kans is dat een groen pixel gras is. Deze kans is namelijk gelijk aan de kans op groen gras, maal de kans op gras, gedeeld door de kans op een groen pixel:

$$P(\text{gras} | \text{groen}) = \frac{P(\text{groen} | \text{gras}) \cdot P(\text{gras})}{P(\text{groen})}$$

$$P(\text{gras} | \text{groen}) = \frac{\frac{3}{4} \cdot \frac{1}{2}}{\frac{5}{8}} = \frac{3}{5}$$

In figuur 5.1 is te zien dat drie vijfde van alles wat groen is inderdaad gras is.



Figuur 5.1. Gebied met $\frac{1}{4}$ huizen, $\frac{1}{4}$ bos en $\frac{1}{2}$ gras (waarvan $\frac{1}{4}$ geel en $\frac{3}{4}$ groen).

Een classificatiemethode op grond van de formule van Bayes geeft de best mogelijke classificatie, onder voorwaarde dat alle typen fouten even ongewenst zijn. Hiervoor moeten natuurlijk de termen in de formule van Bayes nauwkeurig bepaald worden. Hoe deze drie kansen berekend worden, zal hieronder beschreven worden.

De kans op een featurevector $P(\chi_p)$ wordt vaak buiten beschouwing gelaten bij crisp classificaties, omdat deze voor alle klassen gelijk is en dus geen invloed heeft op de beslissing welke de meest waarschijnlijke is. Bij een fuzzy classificatie die kansen als uitvoer geeft is deze kans is echter wel van belang. De kans op een bepaalde featurevector kan niet goed bepaald worden uit het aantal keer dat deze in het beeld voor komt. Omdat er veel meer mogelijke featurevectoren zijn dan er in een beeld aanwezig zijn, is een beeld hiervoor namelijk een te kleine steekproef. De kans op een featurevector kan daarom beter berekend worden met:

$$P(\chi_p) = \sum_{k=1}^K P(\chi_p | C_k) \cdot P(C_k)$$

Door het toepassen van deze formule wordt de som van de a posteriori kansen voor een pixel gelijk aan 1. Daarom wordt dit ook wel normalisatie van de kansen genoemd. Ook bovenstaande formule kan toegelicht worden aan de hand van figuur 5.1.

De a priori kans op een klasse $P(C_k)$ wordt vaak voor alle klassen gelijkgehouden, om de classificatie niet te veel te sturen in de richting van de meest voorkomende klassen. Bij probabilistische segmentatie wordt hier juist wel gebruik van gemaakt. Door deze kans lokaal te schatten kan de classificatie verbeterd worden (zie paragraaf 5.3).

De conditionele kans op een featurevector $P(\chi_p | C_k)$ vormt de kern van een statistische classificatiemethode. Verschillende classificatiemethoden schatten deze kans op verschillende manieren (zie volgende paragraaf).

5.1.2. Conditionele kans op een featurevector

In deze subparagraaf zullen enkele classificatiemethoden beschreven worden, waarmee de kans $P(\chi_p | C_k)$ berekend kan worden. Eerst zullen de maximum-likelihood-classificatie en twee afgeleide methoden behandeld worden, vervolgens de k-nearest-neighbours-classificatie. Er zijn nog veel meer mogelijke methoden en tussenvormen. Voor dit onderzoek zijn echter alleen deze vier methoden overwogen.

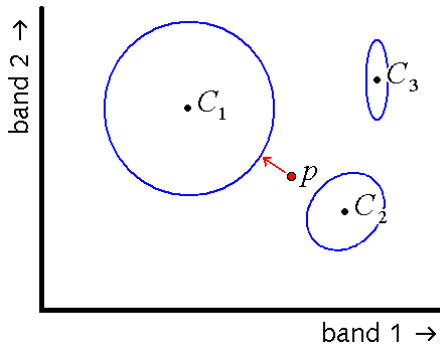
Maximum-likelihood-classificatie

In het trainingsstadium van een maximum-likelihood-classificatie worden uit de trainings-samples de klassengemiddelden en -spreiding bepaald, uitgaande van de normale verdeling. Dit gebeurt met de formules voor het gemiddelde en de empirische (co)variantie (zie subparagraaf 4.2.2). Crisp classificatie vindt plaats op grond van de maximale waarschijnlijkheid (zie figuur 5.2). De conditionele kans op een featurevector kan berekend worden met de formule:

$$P(\chi_p | C_k) = (2\pi)^{J/2} |Q_{\chi_{C_k}}|^{-\frac{1}{2}} e^{-\frac{1}{2}M^2}$$

met: Mahalanobis-afstand $M = \sqrt{(\chi_p - \chi_{C_k})^T \cdot Q_{\chi_{C_k}}^{-1} \cdot (\chi_p - \chi_{C_k})}$

J het aantal banden



Figuur 5.2. Maximum-likelihood-classificatie. De conditionele kans op de featurevector van het pixel p is het grootst voor klasse 1.

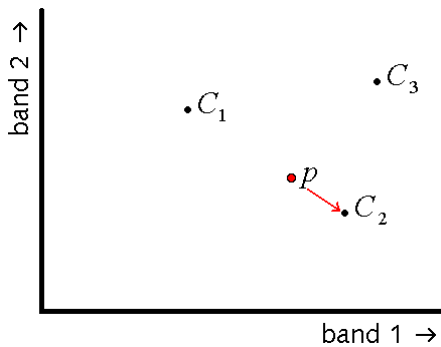
Minimum-distance-classificatie

In het trainingsstadium van minimum-distance-classificatie worden uit de trainingssamples alleen de klassengemiddelden bepaald. Crisp classificatie vindt plaats op grond van de minimale Euclidische afstand in de featurespace (zie figuur 5.3). Minimum-distance-classificatie is een niet-statistische classificatiemethode, maar er kunnen kansen berekend worden uit de Euclidische afstanden met een vereenvoudigde versie van bovenstaande formule. Hiervoor moet $Q_{\chi_{C_k}}$ vervangen worden door een eenheidsmatrix, geschaald met één variantiefactor σ_E^2 voor alle klassen en alle banden:

$$P(\chi_p | C_k) = (2\pi)^{J/2} \sigma_E^{-J} e^{-\frac{1}{2\sigma_E^2} E^2}$$

met: Euclidische afstand $E = \sqrt{(\chi_p - \chi_{C_k})^T \cdot (\chi_p - \chi_{C_k})}$

J het aantal banden



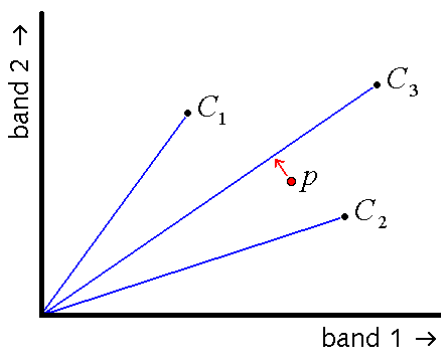
Figuur 5.3. Minimum-distance-classificatie. De conditionele kans op de featurevector van het pixel p is het grootst voor klasse 2.

Dit komt overeen met een maximum-likelihood-classificatie waarbij alle klassen gelijk gespreid zijn, de klassen in elke band dezelfde spreiding hebben en de banden niet gecorreleerd zijn.

Spectral angle mapper

In het trainingsstadium van een classificatie met de spectral angle mapper worden uit de trainingssamples de klassengemiddelden bepaald. Crisp classificatie vindt plaats op grond van de minimale spectrale hoek in de featurespace (zie figuur 5.4). De spectral angle mapper is een vereenvoudiging van minimum-distance-classificatie. Er zijn twee uitgangspunten die aanleiding geven tot deze vereenvoudiging. Ten eerste helpt de intensiteit van een pixel niet veel bij het onderscheiden van verschillende klassen. Ten tweede veroorzaken intensiteitsverschillen binnen een klasse (bijvoorbeeld door ongelijkmatige belichting) veel heterogeniteit. Daarom wordt verondersteld dat het beter is de afstand tot de oorsprong buiten beschouwing te laten bij classificatie.

Ook de spectral angle mapper is een niet-statistische classificatiemethode, maar ook met behulp hiervan kan geprobeerd worden kansen te bepalen. Bij het expertsysteem [Skidmore e.a. 2001] wordt dit bijvoorbeeld gedaan door 1 gedeeld door de hoek te nemen. Het zou wellicht beter zijn om een aangepaste versie van bovenstaande formule te gebruiken.



Figuur 5.4. Spectral angle mapper. De conditionele kans op de featurevector van het pixel p is het grootst voor klasse 3.

k-Nearest-neighbours-classificatie

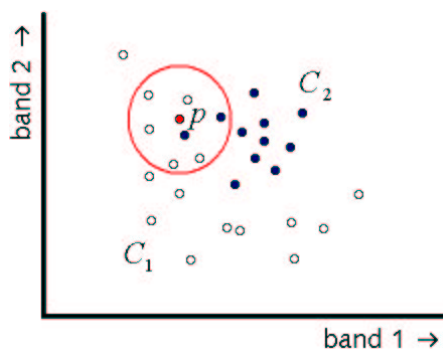
In het trainingsstadium van k-nearest-neighbours-classificatie worden de trainingssamples alleen gelabeld met een klasse. k-Nearest-neighbours-classificatie is namelijk een niet-parametrische classificatiemethode. Dit houdt in dat er geen parameters als klassengemiddelde of (co)varianties geschat worden. Belangrijk gevolg hiervan is dat deze methode voor het schatten van de kansen niet uitgaat van de normale verdeling, zodat ook voor niet normaal verdeelde klassen kansen geschat kunnen worden. Deze methode gaat uit van de dichtstbijzijnde k trainingssamples in de featurespace rond een te classificeren pixel (zie

figuur 5.5). Bij crisp classificatie wordt de klasse toegekend waarvan het label het vaakst voorkomt onder deze k samples. De kans op een klasse is proportioneel met het aantal trainingssamples van een klasse bij deze k nearest neighbours. Wanneer er tijdens het trainingsstadium niet voor alle klassen even veel samples opgegeven zijn, zal de kans op een klasse met meer samples te hoog geschat worden. Hiervoor moet gecorrigeerd worden door te delen door het aantal samples van die klasse:

$$P(\chi_p | C_i) \sim \frac{n_{C_i}}{N_{C_i}} \quad [\text{Gorte 1998 blz. 54}]$$

met: n_{C_i} het aantal trainingssamples voor klasse i onder de k nearest neighbours

N_{C_i} het totaal aantal trainingssamples voor klasse i



Figuur 5.5. k -Nearest-neighbours-classificatie met $k=7$. De conditionele kans op de featurevector van het pixel p is het grootst voor klasse 1.

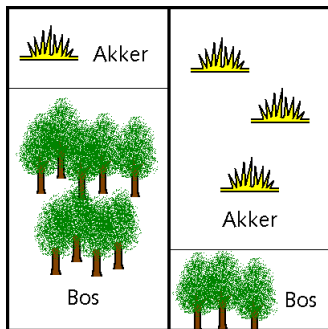
De keuze voor één van de methoden is erg afhankelijk van de toepassing. De ene methode leent zich meer voor een bepaald type data dan een andere methode. De keuze voor een classificatiemethode voor de Schiermonnikoog-data wordt gemaakt in hoofdstuk 6.

5.1.3. Lokale a priori kansen

De a priori kans op een klasse $P(C_k)$ wordt voor classificaties vaak voor alle klassen gelijk gehouden. Enerzijds gebeurt dit omdat deze kansen onbekend zijn; er wordt juist geclassificeerd om te bepalen waar welke klassen aanwezig zijn. Anderzijds wil men de classificatie niet te veel sturen in de richting van de meest voorkomende klassen, omdat anders klassen die weinig voorkomen nauwelijks nog gevonden zouden worden. A priori kansen die voor het gehele beeld gelden geven dan ook slechts een lichte verbetering. Als deze kans echter lokaal bekend is, kunnen wel flinke verbeteringen verkregen worden. Wanneer de klassen namelijk in andere onderlinge verhoudingen in de segmenten voorkomen dan in het totale beeld, zullen de a priori kansen de classificatie gemiddeld wat meer in de goede richting sturen. Het maakt hierbij niet uit of de klassen binnen een segment wel of niet ruimtelijk gescheiden zijn. Dit zou daarom de classificatie van scherp begrensde ruimtelijke objecten met gradueel variërende interne velden in natuurlijke vegetatie wel eens kunnen verbeteren.

Voorbeeld

Wanneer het linker segment in figuur 5.6 geclassificeerd wordt met de a priori kansen $P(C_{Akker}) = 0,25$ en $P(C_{Bos}) = 0,75$ kan een hogere overall-accuracy verwacht worden dan bij gelijke a priori kansen. Hetzelfde geldt voor het rechter segment als $P(C_{Akker}) = 0,75$ en $P(C_{Bos}) = 0,25$ gebruikt worden [Gorte 1998 blz. 31].



Figuur 5.6. Twee segmenten met verschillende klassenaandelen.

Lokale a priori kansen kunnen een oplossing vormen voor spectrale overlap van klassen [Gorte 1998 blz. 43]. De a priori kans helpt dan namelijk bij het maken van de goede keuze. Vooral wanneer de klassen met de spectrale overlap in verschillende segmenten voorkomen zal dit goed werken, maar ook als ze door elkaar in een segment voorkomen zal een lokale a priori kans de classificatie gemiddeld vaker de goede keuze laten maken.

Voorbeeld

Verdroogd gras kan net zo geel zijn als tarwe, waardoor een spectrale overlap bestaat tussen de klassen gras en tarwe. Een geel pixel zal in het algemeen wat vaker tarwe dan gras zijn. Een geel pixel in een groen segment is echter waarschijnlijk gras, terwijl een zelfde geel pixel in een geel segment waarschijnlijk tarwe is.

Op het moment dat er segmenten zouden zijn met een mengsel van gras en tarwe wordt dit moeilijker te bepalen. In agrarische gebieden komt iets dergelijks nauwelijks voor, maar bij natuurlijke vegetatie wel. Als de verhouding van gras en tarwe lokaal bekend is, zullen de a priori kansen er voor zorgen dat de classificatie eerder voor de meest voorkomende klasse kiest. Dit verbetert de classificatie.

Om een classificatie te verbeteren met behulp van lokale a priori kansen, moet er aan twee voorwaarden voldaan worden [Gorte 1998 blz. 31]. Men moet beschikken over:

1. Een opdeling in segmenten zodat verschillende klassenverhoudingen verkregen worden (zie hoofdstuk 6);
2. Schattingen voor de klassenaandelen in de segmenten, die gebruikt kunnen worden als a priori kansen $P(C_k)$ (zie volgende subparagraaf).

5.1.4. Iteratief schatten van lokale kansen

Lokale a priori kansen kunnen de classificatie verbeteren. Probleem is dat deze a priori kansen meestal niet bekend zijn. Het zal over het algemeen geen optie zijn deze gegevens in het veld in te winnen. Dit levert namelijk zo veel veldwerk op dat het voordeel van remote sensing ten opzichte van het karteren middels veldwerk tenietgedaan zou kunnen worden.

Het is echter mogelijk om de lokale a priori kansen te schatten uit een automatische classificatie van het remote-sensing-beeld. Hiervoor is een gewone Bayesiaanse classificatie nodig die als uitvoer de kansen geeft. Zoals in paragraaf 5.2 beschreven is, wordt hierbij de kans $P(\chi_p | C_k)$ geschat. Deze kansen worden, omdat de verhoudingen nog niet bekend zijn, vermenigvuldigd met a priori kansen $P(C_k)$ die voor alle klassen gelijkgehouden worden. Vervolgens wordt er genormaliseerd zodat er a posteriori kansen $P(C_k | \chi_p)$ ontstaan.

Een dergelijke a posteriori kans, bijvoorbeeld $P(C_1 | \chi_p) = 0,04$, betekent dat als er 100 pixels met featurevector χ_p zijn, er waarschijnlijk vier tot klasse 1 zullen behoren. Wanneer we in ieder segment per klasse de som over alle a posteriori kansen berekenen, levert dat een oppervlakteschatting per klasse in dat segment op [Duda en Hart 1973 in Gorte 1998

blz. 33]. Wanneer we deze oppervlakteschattingen delen door de segmentgrootte worden klassenaandelen voor dat segment verkregen. Deze klassenaandelen kunnen gebruikt worden als lokale a priori kansen. Voor grote segmenten zullen de aandelen nauwkeuriger zijn dan voor kleine segmenten, omdat ze op basis van meer pixels bepaald zijn.

De formule voor de a priori kans op een klasse in een segment is:

$$P(C_k) = \frac{1}{N} \cdot \sum_{p=1}^N P(C_k | \chi_p)$$

met: N het aantal pixels in het segment.

De zo verkregen a priori kansen zullen een fout bevatten doordat er voor het bepalen hiervan uitgegaan is van gelijke a priori kansen op alle klassen. Het zal echter beter zijn om deze te gebruiken dan te veronderstellen dat alle klassen in alle segmenten gelijke aandelen hebben. Dit geeft de mogelijkheid om de lokale a priori kansen iteratief te schatten. De voorwaarde hiervoor is dat wanneer de gebruikte a priori kansen een benadering zijn, de geschatte a priori kansen nauwkeuriger zijn dan de eerste benadering. Dat aan deze voorwaarde voldaan wordt is bewezen door Gorte [1998 blz. 36-42] voor de schatting van niet-parametrische a priori kansen. De iteratie procedure ziet er als volgt uit:

Iteratieprocedure

1. Initialiseer de a priori kansen (bijvoorbeeld allemaal $1/K$).
2. Pas de formule van Bayes toe: Vermenigvuldigen met de a priori kansen; Normaliseren van de kansen.
3. Neem voor iedere klasse de som over de pixels per segment. Deel door de segmentgrootte. Dit geeft verbeterde a priori kansen.
4. Ga naar 2.

De geschatte lokale a priori kansen kunnen gebruikt worden om de formule van Bayes toe te passen, zodat een verbeterde pixelgewijze statistische fuzzy classificatie verkregen wordt. Natuurlijk kunnen nog fouten gemaakt worden, soms zelfs op plaatsen waar dat zonder lokale a priori kansen wel goed ging. Over het gehele beeld zal de classificatie er echter beter van worden, mits de kansen $P(\chi_p | C_k)$, waaruit de a priori kansen bepaald zijn, kloppen.

5.2. Segmentatie

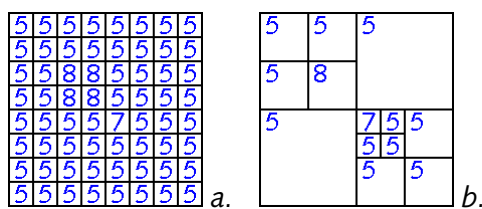
De segmenten, waar de lokale a priori kansen voor geschat worden, kunnen op verschillende manieren verkregen worden, bijvoorbeeld uit additionele gegevens. Zelfs een classificatie met postcodegebieden als segmenten kan succesvol zijn [Gorte 1998 blz. 51]. Omdat additionele gegevens niet altijd beschikbaar zijn is het wenselijk dat ook de opdeling uit de beelden zelf gehaald wordt door middel van beeldsegmentatie. Doel van zo'n segmentatie is het beeld te verdelen in aaneengesloten groepen pixels op grond van hun spectrale kenmerken. Het is hierbij de bedoeling dat de segmenten zoveel mogelijk corresponderen met terreinobjecten. Terreinobjecten zullen namelijk over het algemeen bestaan uit één of enkele klassen, waardoor het gebruik van a priori kansen het meest effectief is.

5.2.1. Segmentatiepiramide

De methode voor beeldsegmentatie gaat uit van het samenvoegen van regio's (regio-merging) gebaseerd op zogenaamde quadrees. Het blijkt dat dit voor agrarische gebieden goede resultaten en relatief korte rekentijden geeft [Gorte 1998 blz. 83].

Een quadtree is een methode om rasterdata op te slaan (zie figuur 5.7). Hierbij worden zogenaamde leaves onderscheiden die een waarde hebben. Het verschil met pixels is dat de leaves verschillende groottes kunnen hebben. De grootte van een leaf wordt bepaald door het niveau. Een leaf in het laagste niveau is zo groot als één pixel, in ieder hoger niveau beslaat een leaf vier leaves uit het niveau daaronder. Een quadtree bestaat uit twee bestanden: een qtl-bestand met daarin het niveau van de leaves en een qtv-bestand met daarin de waarden van de leaves.

Bij quadrees hoeft in gebieden met dezelfde pixelwaarde deze waarde niet voor elk pixel afzonderlijk opgeslagen te worden. Het grote voordeel van quadrees is echter niet zozeer de datacompressie, maar het gebruik van programma's die direct met quadrees werken. Programma's die quadrees gebruiken zonder deze eerst te converteren naar een rasterbestand kunnen aanzienlijk sneller zijn dan programma's die met rasterbestanden werken. Daarom worden voor de segmentatie de beelden naar quadrees geconverteerd, en pas aan het eind van de probabilistische segmentatieprocedure worden de resultaten weer geconverteerd naar rasterbestanden.



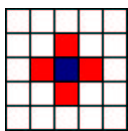
Het programma gebruikt drie banden voor de segmentatie. De segmentatiemethode voegt op grond van deze drie banden recursief aangrenzende leaves en regio's samen tot regio's indien:

1. De Euclidische afstand in de featurespace tussen twee kandidaten kleiner is dan een drempelwaarde d_1 .
2. De variantie van de pixelwaarden in elke band van het samengevoegde segment, alsmede de covarianties tussen de banden in dat segment, niet groter worden dan een drempelwaarde d_2 .

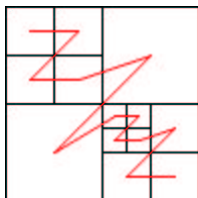
Deze twee drempelwaarden worden samen uit één drempelwaarde d bepaald. Omdat (co)varianties kwadratisch zijn is de drempel d_2 gedefinieerd als het kwadraat van d . De verhouding tussen d_1 en d_2 is enigszins willekeurig. In het gebruik is gebleken dat een factor 3 in onderstaande formule redelijk goede resultaten geeft.

$$d_1 = 3 \cdot d \quad \text{en} \quad d_2 = d^2$$

Leaves en regio's grenzen aan elkaar wanneer ze een gemeenschappelijke rand hebben. Ieder pixel heeft daarmee vier burens (zie figuur 5.8). De Z-scan-volgorde (zie figuur 5.9) waarin geprobeerd wordt om leaves en regio's samen te voegen, is van invloed op het resultaat. Hierdoor ontstaan hoekig gevormde segmenten. Om dit te voorkomen kan er iteratief gesegmenteerd worden met een geleidelijk toenemende drempelwaarde. In de eerste stap worden leaves samengevoegd volgens bovenstaande regels met een lage drempelwaarde. Vervolgens worden in de tweede stap met een net iets hogere drempelwaarde de in stap één ontstane segmenten samengevoegd, volgens dezelfde bovenstaande regels. Bij iedere volgende stap wordt de drempelwaarde verder verhoogd.

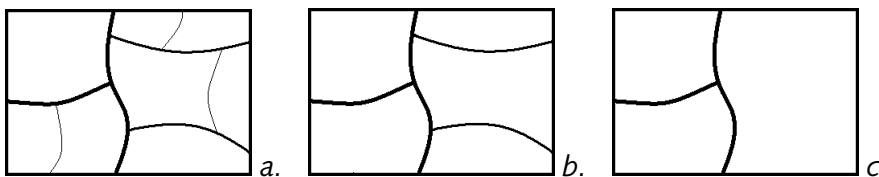


Figuur 5.8. De vier aangrenzende pixels van een pixel.



Figuur 5.9. Voorbeeld van de Z-scan-volgorde [Gorte 1998 blz. 72].

Het iteratief segmenteren met oplopende drempelwaarde levert een segmentatiepiramide op, waarbij iedere drempelwaarde een niveau van de piramide geeft. (zie figuur 5.10). In de hogere lagen van de piramide zijn er minder segmenten en zijn de segmenten groter, omdat het aggregaties zijn van de kleinere segmenten in een lager niveau. Elk segment heeft daarom in ieder hoger niveau één supersegment, terwijl het in ieder lager niveau meerdere subsegmenten kan hebben. Uit deze segmentatiepiramide kan een gewenst niveau gekozen kan worden.



Figuur 5.10. Segmentatiepiramide met drie niveaus a, b en c [Gorte 1998 blz. 87].

Bij de segmentatie ontstaan op de grenzen van grote segmenten soms kleine segmenten van één pixel. Dit zijn mixed pixels die te verschillend zijn om bij één van de twee grotere segmenten toegevoegd te worden. Nadeel van deze kleine segmenten is dat ze onnauwkeurige a priori kansen zullen geven. Het is echter de vraag of het wenselijk is om dergelijke pixels aan een groot segment toe te voegen. Dit zou dat segment alleen maar vervuilen. De mate waarin kleine segmenten aan een groot segment worden toegevoegd is afhankelijk van de verhouding d_1 en d_2 . In bepaalde situaties kan daarom een wat hogere of lagere factor dan 3 in bovenstaande formule geprefereerd worden.

5.2.2. Segmentsselectie

Een probleem is dat het onduidelijk is welk niveau uit de piramide de beste segmentatie oplevert. En als een redelijk niveau gevonden is, zullen in sommige gebieden in het beeld de segmenten te groot kunnen zijn, terwijl in andere delen juist nog te kleine segmenten aanwezig zijn. Een segment waarin bijvoorbeeld twee klassen voorkomen is te groot als het opgesplitst zou kunnen worden in twee segmenten met allebei maar één klasse. Een segment van maar een paar pixels is te klein als aangrenzende segmenten dezelfde klassen bevatten. De oplossing hiervoor is het selecteren van segmenten uit verschillende niveaus van de piramide, op basis van de klassenaandelen. Deze klassenaandelen worden voor de gehele segmentatiepiramide bepaald (zoals beschreven is in paragraaf 5.3). Gorte [1998 blz. 93] selecteert alle pure segmenten S waarvoor geldt:

$$S : (n(S) = 1) \wedge \left(\forall_{T \in \text{super}(S)} : n(T) \neq 1 \right)$$

en alle gemengde segmenten waarvoor geldt:

$$S : (n(S) \neq 1) \wedge \left(\forall_{T \in \text{super}(S)} : n(T) \neq 1 \right) \wedge \left(\forall_{T \in \text{sub}(S)} : n(T) \neq 1 \right) \wedge \left(\forall_{T \in \text{super}(S)} \exists_{U \in \text{sub}(T)} : n(U) = 1 \right)$$

met: $\text{super}(S)$ supersegmenten van S in alle bovenliggende niveaus;
 $\text{sub}(S)$ subsegmenten van S in alle onderliggende niveaus;
 $n(S)$ aantal klassen in segment S .

Dit komt er op neer dat, indien mogelijk, pure segmenten geselecteerd worden. Alleen waar dat niet mogelijk is worden heterogene segmenten geselecteerd. In beide gevallen gaat de voorkeur uit naar zo groot mogelijke segmenten.

5.2.3. Nieuwe methode voor segmentselectie

Hoewel dit goed werkt voor agrarische gebieden, werd bij dit onderzoek verwacht dat dit voor natuurlijke vegetatie misschien niet zo succesvol zou zijn. Er zouden namelijk wel eens weinig pure segmenten gevonden kunnen worden, met als resultaat hele grote heterogene segmenten. Zulke grote segmenten zullen net als a priori kansen voor het gehele beeld slechts weinig verbetering geven. Omdat de selectie niet in een operationeel programma verwerkt was, is daarom een andere selectieregel geprogrammeerd. Hierbij worden alle segmenten S geselecteerd waarvoor geldt:

$$S : \left(\forall_{T \in \text{sub}(S)} : n(T) \geq n(S) \right) \wedge \left(\forall_{T \in \text{super}(S)} \exists_{U \in \text{sub}(T)} : n(U) < n(T) \right)$$

Deze formule controleert in twee stappen of een segment geselecteerd moet worden. Het eerste deel van de formule zorgt er voor dat er geen subsegmenten zijn met minder klassen. In dat geval zouden die namelijk geselecteerd moeten worden. In de tweede stap wordt gekeken of er geen supersegmenten zijn die beter geselecteerd zouden kunnen worden. Een supersegment is beter als dit supersegment evenveel of minder klassen heeft dan al zijn eigen subsegmenten.

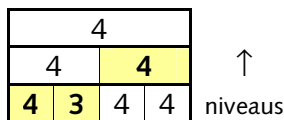
Dit komt er op neer dat de segmenten met zo min mogelijk klassen geselecteerd worden. Als er meerdere kandidaten zijn worden zo groot mogelijke segmenten gekozen (zie figuur 5.11).

1	3				2
1	2	3	2	2	
1	1	2	2	2	2
					2

↑
niveaus

Figuur 5.11. Schematische weergave van een segmentatiepiramide met in ieder segment het aantal aanwezige klassen. De geselecteerde segmenten zijn gemarkeerd

In één bepaalde situatie neemt deze selectiemethode een vreemde beslissing. Een groot segment met maar een paar klassen, dat in één van de subsegmenten één klasse minder heeft, zal niet geselecteerd worden. In plaats daarvan zullen meerdere kleinere segmenten geselecteerd worden (zie figuur 5.12). Wanneer dit veel erg kleine segmentjes zouden zijn, is dit eigenlijk onwenselijk, omdat dit onnauwkeurige a priori kansen op zal leveren. Dit kan enigszins tegengegaan worden door de functie $n(S)$ iets aan te passen. Het programma biedt namelijk de mogelijkheid om kleine aandelen van een klasse (uitgedrukt in een fractie van het segment of een aantal pixels) te verwaarlozen. Hiermee kan de grootte van de geselecteerde segmenten iets beïnvloed worden. Deze verwaarlozing is alleen van invloed op het selecteren van de segmenten. Voor het bepalen van de lokale a priori kansen vindt geen verwaarlozing plaats.



Figuur 5.12. Schematische weergave van een segmentatiepiramide met in ieder segment het aantal aanwezige klassen. De geselecteerde segmenten zijn gemarkeerd.

5.3. Beschouwing en aanbevelingen

Bij de probabilistische segmentatie wordt uitgegaan van een fuzzy statistische classificatie, die op grond van alle banden de kansen op de verschillende klassen berekent. Daarnaast wordt een segmentatiepiramide gemaakt op grond van drie banden. Vervolgens worden voor de gehele piramide de klassenaandelen per segment bepaald. Uit verschillende lagen van de piramide worden segmenten geselecteerd met zo min mogelijk klassen. De klassenaandelen van de geselecteerde segmenten kunnen gebruikt worden als a priori kansen, om op grond van de formules van Bayes de eerste classificatie te verbeteren (zie figuur 5.14).

Eigenlijk worden er drie fuzzy classificaties gemaakt:

- Fuzzy classificatie per pixel (de kans op iedere klasse);
- Fuzzy classificatie per segment (een oppervlakteaandeel per klasse);
- Verbeterde fuzzy classificatie per pixel (de kans op iedere klasse);

Door de maximale kans (of het grootste oppervlakteaandeel) te selecteren, worden drie crisp classificaties verkregen. De verbeterde fuzzy en crisp classificaties zijn de beste classificaties. Voor monitoring zou echter ook de voorkeur gegeven kunnen worden aan een van de andere classificaties.

Doordat de segmentsgewijze fuzzy classificatie bepaald is op basis van meerdere pixels is deze minder gevoelig voor meetruis dan een pixelgewijze fuzzy classificatie. Voorwaarde hiervoor is natuurlijk dat de segmenten niet te klein zijn. Deze aandelen per segment lenen zich daarom wellicht beter voor monitoring van geleidelijke veranderingen dan een pixelgewijze classificatie. Om te voorkomen dat voornamelijk verschillen ten gevolge van een andere segmentatie te zien zijn, zou besloten kunnen worden om voor alle beelden dezelfde segmentatie te gebruiken.

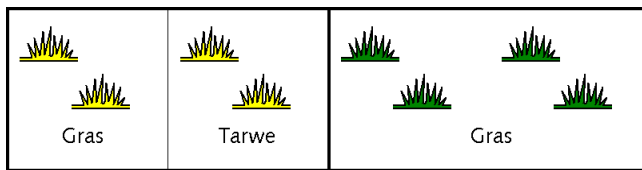
Eén van de voordelen van probabilistische segmentatie is dat spectrale overlap van klassen een minder groot probleem is dan bij een pixelgewijze classificatie. Dat wil echter niet zeggen dat dit alle problemen oplost. Bij spectrale overlap zal namelijk ook het schatten van de a priori kansen moeilijker zijn. Wanneer bij het schatten van de a priori kansen op de klassen in een segment een fout gemaakt wordt, zal dit grote gevolgen hebben voor de classificatie van de pixels binnen dat segment.

De verschillende kansen in de formules van Bayes moeten op hetzelfde gebied van toepassing zijn. De kansen $P(\chi_p | C_k)$ en $P(C_k)$ zijn dit echter niet. De kans $P(C_k)$ wordt lokaal bepaald terwijl $P(\chi_p | C_k)$ uit de trainingsset bepaald wordt voor het gehele beeld. Wanneer deze kans toch gebruikt wordt in combinatie met de lokale $P(C_k)$ wordt impliciet verondersteld dat de waarden die voor het gehele beeld gelden ook lokaal van toepassing zijn. Dit is niet altijd terecht (zie voorbeeld).

Voorbeeld

De kans op een geel pixel, gegeven dat het een graspixel is, zal in het algemeen (voor het gehele beeld) klein zijn. Lokaal kan deze kans daar echter van afwijken. In een veld dat een andere bodemsoort heeft die sneller tot verdroging leidt, waardoor dit hele veld (bijna) geheel uit geel gras bestaat, is deze kans veel groter. In het linker

segment in figuur 5.13 is de kans $P(\text{Geel} | \text{Gras}) = 1$ terwijl deze kans voor het gehele gebied $\frac{1}{3}$ is.



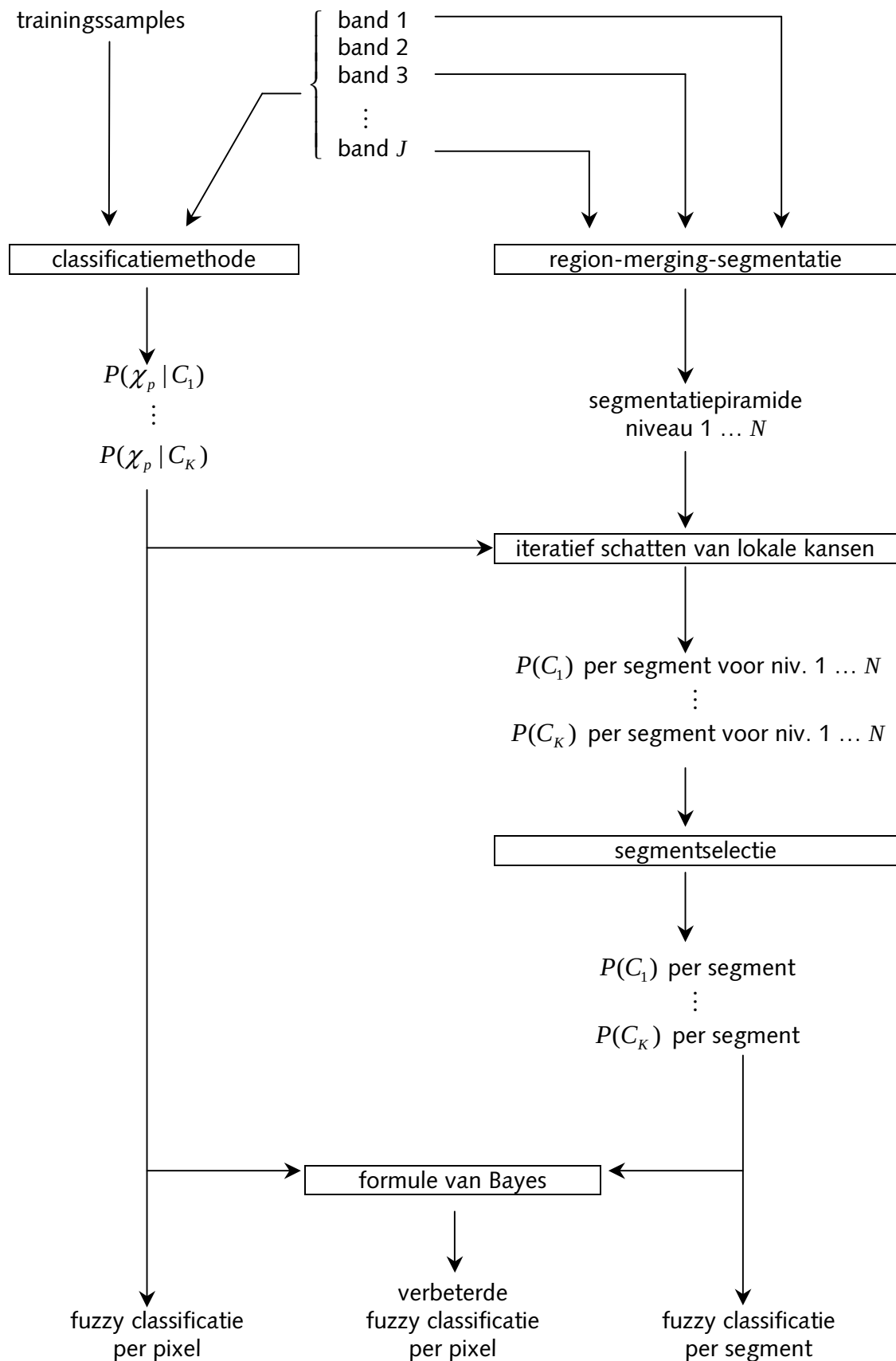
Figuur 5.13. Gebied met twee segmenten. Een geel segment met gras en tarwe en een groen segment met alleen gras.

Gorte [1998 blz. 55-56] beschrijft een methode waarmee ook de conditionele kans op een featurevector lokaal geschat kan worden. Deze methode is echter alleen toepasbaar wanneer men over erg veel trainingssamples beschikt. Omdat dat bij dit onderzoek niet het geval is, wordt er van uitgegaan dat de kansen die voor het gehele beeld bepaald zijn ook lokaal van toepassing zijn.

De rekentijd van het programma kan oplopen tot 11 uur. Dit zou nog iets verbeterd kunnen worden door het programma efficiënter te maken, bijvoorbeeld door een toets op significantie bij de iteratieve schatting van de a priori kansen toe te voegen. Ook zou het schrijven van C-programma's in plaats van de bash-scripts een verbetering kunnen geven.

Het programma gebruikt drie banden voor de segmentatie. Bij data met meer banden is het wenselijk dat de informatie uit alle banden gebruikt wordt. Het programma zou hiervoor aangepast moeten worden. Dit is uit tijdsoverweging niet gedaan. De regels voor het samenvoegen van leaves en regio's zouden dan ook gewijzigd moeten worden. In ieder geval moet de verhouding tussen de twee drempelwaarden dan aangepast worden aan de toenemende spectrale afstanden in een meerdimensionale featurespace. Wanneer de drie te gebruiken banden goed gekozen worden, of wanneer bijvoorbeeld de eerste drie componenten van een principal-components-analyse (PCA) worden genomen, geeft dit waarschijnlijk een net zo goede segmentatie.

Afgezien hiervan is de enige noodzakelijke aanpassing, voor het gebruik van hyperspectrale beelden voor probabilistische segmentatie, het bepalen van kansen op de verschillende kansen. Door het relatief grote aantal banden ten opzichte van het aantal samples is het moeilijk deze kansen te schatten (zie paragraaf 6.1). Dit probleem ondervindt iedere statistische classificatiemethode die op hyperspectrale beelden toegepast wordt.



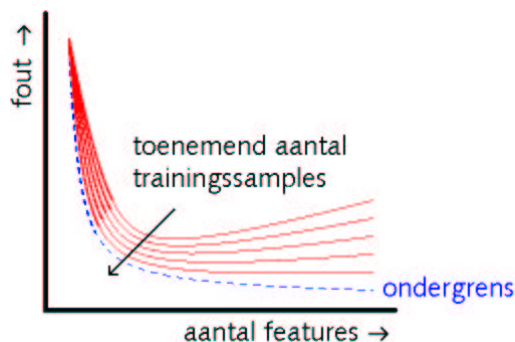
Figuur 5.14. Schema van de probabilistische segmentatie.

6. Conditionele kans op een featurevector

In het vorige hoofdstuk (subparagraaf 5.1.2) zijn enkele methoden besproken voor het schatten van de conditionele kans op een featurevector $P(\mathcal{X}_p | C_k)$. In dit hoofdstuk zal de keuze voor één van deze schattingsmethoden gemotiveerd worden. Van belang is dat de methode realistische kansen schat voor de gebruikte hyperspectrale beelden en een beperkte hoeveelheid veldgegevens. Daarom moet deze niet te gevoelig zijn voor de overtraining

6.1. Overtraining

Veel classificatiemethoden, waaronder maximum-likelihood-classificatie en k-nearest-neighbours-classificatie, kunnen last hebben van de zogenaamde curse of dimensionality [Jain en Waller 1978, Duin 1978 en van Otterloo en Young 1978 in Hoekstra 1998 paragraaf 1.2]. Dit is het verschijnsel dat de classificatie slechter presteert door het toevoegen van features (banden) bij gelijk een aantal trainingssamples (zie figuur 6.1). Bij het gebruik van hyperspectrale data kan dit zich al snel voordoen.

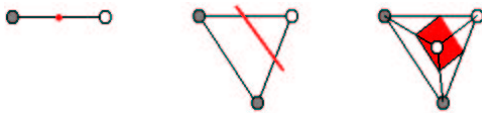


Figuur 6.1. Curse of dimensionality. Voor een vast aantal trainingssamples zal de classificatiefout van een classificatiemethode uiteindelijk toenemen wanneer het aantal features wordt verhoogd [Hoekstra 1998 paragraaf 1.2].

De verslechtering wordt veroorzaakt doordat bij het gebruik van veel features de classificatiemethode zich meer gaat aanpassen aan de ruis in de trainingsset. Bij validatie met de trainingsset lijkt het toevoegen van features dan verbetering te geven, maar bij validatie met een onafhankelijke testset blijkt dit niet het geval te zijn. Dat het toevoegen van extra features (en dus informatie voor de training) soms toch een verslechtering geeft, is de reden dat dit ook wel overtraining genoemd wordt. Het onderstaande voorbeeld moet het verschijnsel enigszins verhelderen.

Voorbeeld

We beschouwen de situatie met twee klassen, J features en $J+1$ trainingssamples. Iedere classificatiemethode levert uiteindelijk een beslissingsgrens die de J -dimensionale featurespace opdeelt. De eenvoudigste beslissingsgrens is een lineaire deelruimte met dimensie $J-1$ (zie figuur 6.2). In het geval $J+1$ trainingssamples kan zelfs een classificatiemethode met zo'n eenvoudige grens deze zo kiezen dat de trainingssamples perfect in de twee klassen opgedeeld worden. Validatie op de trainingsdata zou een overall-accuracy van 100% geven, de classificatiemethode is volledig aangepast aan de ruis in de trainingsset.



Figuur 6.2. In een 1-dimensionale featurespace kunnen 2 samples gescheiden worden door een punt. In een 2-dimensionale featurespace kunnen 3 samples gescheiden worden door een lijn. In een 3-dimensionale featurespace kunnen 4 samples gescheiden worden door een vlak.

Het blijkt dat een classificatiemethode dus enigszins moet generaliseren om goede resultaten op te leveren. Wanneer de overvullingheid te klein wordt worden de resultaten slechter.

Een remedie tegen overtraining is het verhogen van het aantal samples, maar vaak is dat niet mogelijk. Dit zou namelijk zeer omvangrijk veldwerk vereisen. Een andere optie is een andere classificatiemethode te kiezen die minder gevoelig voor overtraining is. In het geval van parametrische classificatiemethoden zijn dit bijvoorbeeld methoden die minder parameters schatten. Een derde mogelijke oplossing is feature-reductie. Door gebruik te maken van minder banden kan overtraining voorkomen worden en worden de resultaten beter. Dit lijkt vreemd want de informatie in de niet gebruikte banden wordt weggegooid. De kunst is dan ook om het aantal banden zo te reduceren dat er zo veel mogelijk van de informatie behouden wordt. Er zijn twee verschillende strategieën voor feature-reductie, namelijk feature-selectie en feature-extractie. Bij selectie wordt een deel van de banden geselecteerd voor de classificatie, de overigen worden niet gebruikt. Bij extractie worden nieuwe banden gevormd met daarin zoveel mogelijk informatie van alle banden.

Enkele mogelijke methoden voor feature-reductie:

- Principal-components-analyse (PCA);
- Het gemiddelde van meerdere banden nemen;
- Selectie van banden op grond van expertkennis;
- Selectie van banden op grond van statistisch bepaalde informatie-inhoud van de banden.

6.2. Keuze voor een schattingsmethode

Het enige criterium voor een geschikte methode voor het schatten van de conditionele kans op een featurevector is dat deze realistische kansen moet schatten. Dit moet blijken uit verbeterde classificatieresultaten bij toepassing van probabilistische segmentatie. In het geval van de gebruikte data is het van belang dat de schattingsmethode geschikt is voor hyperspectrale beelden en maar een beperkt aantal trainingssamples per klasse. Daarom moet deze niet te gevoelig zijn voor overtraining.

In eerste instantie is er van uitgegaan dat maximum-likelihood-classificatie gebruikt zou worden voor het schatten van de kansen. Een probleem dat dan opgelost zou moeten worden is hoe er voor hyperspectrale data getraind kan worden. Voor maximum-likelihood-classificatie zijn namelijk tussen de 10 en 100 trainingssamples per band per klasse benodigd [Gorte 1998 blz. 11]. Ondanks het uitvoerige veldwerk zijn de 208 trainingssamples hoogstens voldoende voor het beschrijven van een stuk of 9 klassen in een paar banden. Hoewel crisp maximum-likelihood-classificatie met klassenindeling niveau 1 (zie paragraaf 3.4) op 3 banden nog redelijk goede resultaten gaf, bleek het toepassen van probabilistische segmentatie met de berekende kansen de resultaten te verslechteren in plaats van te verbeteren. Dit duidt er op dat de berekende kansen niet juist zijn. Dit kan veroorzaakt worden door een tekort aan trainingssamples of doordat de klassen niet normaal verdeeld

zijn. Dit laatste zou overigens niet verwonderlijk zijn omdat de 9 klassen samengesteld zijn uit meerdere vegetatietypen. Samengestelde klassen zijn meestal niet normaal verdeeld, zelfs niet als de samenstellende delen dat wel zijn.

Het onderscheiden van minder dan 9 klassen levert niet voldoende ecologische informatie meer op. Bij het gebruik van minder dan drie banden bestaat er zeker onvoldoende spectraal onderscheid tussen de klassen. Het verwerven van meer trainingssamples is ook niet mogelijk. De enige optie om overtraining te voorkomen die dan nog overblijft is het kiezen van een andere classificatiemethode die bijvoorbeeld een kleiner aantal parameters bepaalt voor de beschrijving van de klassen.

De statistische classificatie op grond van minimum-distance-classificatie (zoals in subparagraaf 5.1.2 is beschreven) schat slechts een gemiddelde per klasse en één variantie in alle banden voor alle klassen. Hierdoor zijn aanzienlijk minder samples nodig. Enkele experimenten met minimum-distance-classificatie gaven echter nog slechtere resultaten dan maximum-likelihood-classificatie. Blijkbaar is de schaarste van samples niet de beperkende factor, maar is de veronderstelling van gelijke spreiding voor alle banden en klassen niet juist. Opmerkelijk was dat er wel vrij redelijke crisp classificaties verkregen werden wanneer er eerst veel klassen werden onderscheiden en deze na de classificatie werden samengevoegd.

Aangezien de spectral angle mapper nog weer een verdere simplificatie van maximum-likelihood-classificatie is dan minimum-distance-classificatie, kan verwacht worden dat deze geen verbetering zal geven. Bij het expertsysteem [Skidmore e.a. 2001] is er echter met succes gebruik gemaakt van de spectral angle mapper. Nu blijkt het zo te zijn dat hierbij geen spectral angle mapper met klassengemiddelden is uitgevoerd, maar met ieder trainingssample als aparte klasse. Na de classificatie worden al deze klassen dan samengevoegd tot de gewenste klassenindeling. Dit blijkt betere resultaten te geven dan het gebruik van klassengemiddelden.

Voor maximum-likelihood-classificatie met meer dan 9 klassen zijn voor de Schiermonnikoog-data dus te weinig samples beschikbaar. De 9 klassen van het eerste niveau van de klassenindeling (zie paragraaf 3.4) zijn niet normaal verdeeld of er zijn meer samples nodig voor maximum-likelihood-training. De minimum-distance-classificatie en de spectral angle mapper gebruiken te weinig parameters om de klassen goed te beschrijven. Wanneer echter eerst veel klassen onderscheiden worden die na de classificatie samengevoegd worden tot een klassenindeling met minder klassen, worden betere resultaten verkregen. Deze wijze van classificeren is een soort k-nearest-neighbours-classificatie met $k=1$. Een belangrijk nadeel van k gelijk aan 1 kiezen is dat een dergelijke classificatiemethode erg gevoelig is voor foutieve trainingssamples. Bovendien kunnen er bij $k=1$ alleen kansen berekend worden als ook nog eens de afstand of hoek in ogenschouw wordt genomen, waardoor een soort tussenvorm ontstaat. Omdat Gorte [1998] met succes een k-nearest-neighbours-classificatie toegepast heeft is besloten om ook bij dit onderzoek deze classificatiemethode te gebruiken voor het schatten van de conditionele kans op een featurevector.

6.3. k-Nearest-neighbours-classificatie

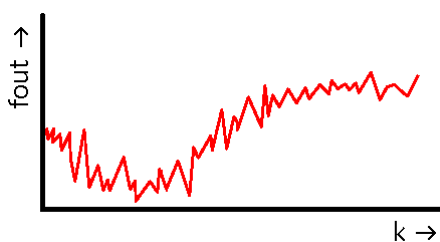
Er zijn verschillende manieren om een k-nearest-neighbours-classificatie te implementeren. Eén methode is een algoritme in de featurespace te laten zoeken naar trainingssamples in de nabijheid van het te classificeren pixel. Hierbij wordt geleidelijk de straal van het te doorzoeken gebied opgevoerd totdat k samples gevonden zijn. Deze methode is gebruikt door Gorte [1998] en is zeer geschikt voor toepassingen met veel trainingssamples en weinig dimensies. Bij veel dimensies wordt deze methode echter zeer rekenintensief. Daarom is bij dit

onderzoek gekozen voor een programma met een andere implementatie. Dit programma berekent voor een te classificeren pixel de afstanden in de featurespace naar alle trainingssamples, om vervolgens de dichtstbijzijnde k samples te selecteren. Uiteraard werkt dit alleen goed bij een niet te groot aantal trainingssamples, maar een groot aantal dimensies geeft geen onacceptabel lange rektijden. Het gebruikte programma is onderdeel van de ANN-programmabibliotheek (Approximate Nearest Neighbour). Deze programmabibliotheek is ontwikkeld bij de Universiteit van Maryland [Mount 1998]. Het programma biedt de mogelijkheid om de rektijd te verkorten door het slechts een benadering (approximate) te laten geven. Deze optie is echter niet gebruikt.

Net zoals andere classificatiemethoden heeft de k -nearest-neighbours-classificatiemethode, zoveel mogelijk samples nodig. Overtraining kan namelijk ook bij k -nearest-neighbours-classificatie optreden [Hoekstra 1998 paragraaf 1.2]. In het kader van dit onderzoek is niet uitgezocht in welke mate zich dit bij de Schiermonnikoog-data voordoet. k -Nearest-neighbours-classificatie met ANN gaf echter ondanks het beperkte aantal trainingssamples redelijke resultaten. De kansen waren goed genoeg om probabilistische segmentatie de resultaten te laten verbeteren (zie hoofdstuk 7).

6.3.1. Keuze voor k

De classificatiefout van een k -nearest-neighbours-classificatie is afhankelijk van de keuze voor k (zie figuur 6.3). De classificatiefout blijkt bij het verhogen van k eerst af te nemen tot een optimum, om vervolgens weer toe te nemen richting een maximale waarde. Classificatie met k gelijk aan N (het totaal aantal trainingssamples) is gelijk aan classificatie volgens de a priori kansen. Als vuistregel voor de optimale keuze geldt: $k \approx \sqrt{N}$ [Duin 2001].



Figuur 6.3. Typisch gedrag van de k -nearest-neighbours-classificatiefout bij verschillende keuzes voor k [Duin 2001].

Het is duidelijk dat k niet groter gekozen moet worden dan het aantal trainingssamples van de klassen met de minste trainingssamples. Bij het schatten van kansen moet k zeker ook niet te klein gekozen worden, daar de mate van detail waarmee de kansen bepaald kunnen worden afhankelijk van k is. Voor crisp classificatie is de keuze $k=1$ wel mogelijk, met als nadeel dat deze de classificatie erg gevoelig maakt voor foutieve trainingssamples. De exacte keuze van k blijkt niet zo erg van invloed op het resultaat te zijn [Gorte 1998 blz. 19].

6.3.2. Fuzzy training voor k -nearest-neighbours-classificatie

In hoofdstuk 4 is een nieuwe fuzzy trainingmethode voor maximum-likelihood-classificatie geïntroduceerd. Het is ook mogelijk om de k -nearest-neighbours-classificatiemethode zo aan te passen dat deze gebruik kan maken van fuzzy trainingssamples. Omdat inmiddels gebleken was dat fuzzy training voor de beschikbare veldwerkdata niet zo zinnig was, is dit niet meer uitgevoerd. Door de geringe precisie van de geometrische transformatie van het beeld is het namelijk onmogelijk exact het juiste pixel in het beeld te identificeren. Hierdoor zullen de klassenaandelen die uit het veldwerk bepaald worden niet overeenkomen met de klassenaandelen in het trainingssample. Doordat bovendien naar zo puur mogelijke samples in het

terrein gezocht was, valt er van fuzzy training met deze data niet veel verbetering te verwachten.

Gezien de voordelen die het gebruik van fuzzy trainingssamples met zich mee zou kunnen brengen zou dit een onderwerp voor verder onderzoek kunnen zijn. Te denken valt aan een methode waarbij net als bij gewone k-nearest-neighbours-classificatie (zie subparagraaf 5.1.2) voor een te classificeren pixel de dichtstbijzijnde k trainingssamples in de featurespace geselecteerd worden. Van deze samples moeten de aandelen van de verschillende klassen bekend zijn, doordat ze bijvoorbeeld tijdens veldwerk bepaald zijn. Vervolgens worden de klassenaandelen over de k trainingssamples per klasse opgeteld. Dan wordt per klasse een waarde verkregen, vergelijkbaar met het aantal neighbours bij gewone k-nearest-neighbours-classificatie. Deze waarden kunnen op de normale manier gebruikt worden voor het bepalen van kansen. Dit levert een fuzzy classificatie per pixel op, verkregen uit fuzzy trainingssamples. Onderzocht zou moeten worden of dit goede schattingen van de kansen op de klassen zijn.

6.4. Beschouwing en aanbevelingen

Bij het gebruik van hyperspectrale data ontstaat het gevaar op overtraining. De beste methode voor het schatten van de conditionele kans op een featurevector blijkt k-nearest-neighbours-classificatie te zijn. Deze geeft voor de Schiermonnikoog-data redelijke classificatieresultaten, die goed genoeg zijn om probabilistische segmentatie de classificatie te doen verbeteren.

Gezien de mogelijke voordelen van het gebruik van fuzzy trainingssamples verdient het aanbeveling om verder onderzoek te doen naar het gebruik van fuzzy trainingssamples voor k-nearest-neighbours-classificatie.

Het gebruikte programma voor k-nearest-neighbours-classificatie is één van de onderdelen van de ANN-programmabibliotheek. Dit programma heeft tekstbestanden nodig voor de invoer. Daarom zijn voor dit onderzoek de rasterbestanden van de banden en trainingssamples geconverteerd van bytebestanden naar tekstbestanden. Door de grote data-omvang van de beelden kost dit veel tijd. Het programma zou aanzienlijk versneld kunnen worden wanneer met behulp van de bibliotheek een programma geschreven zou worden dat rechtstreeks bytebestanden in kan lezen.

7. Evaluatie

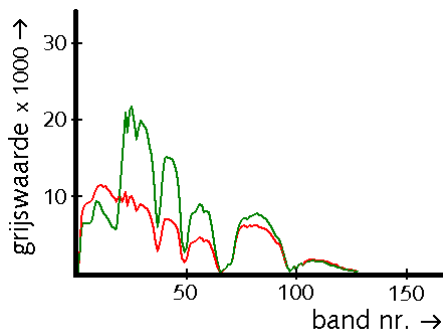
In dit hoofdstuk zal de probabilistische segmentatie op de Schiermonnikoog-data getest worden. Op basis van de validatie van de resultaten zullen conclusies getrokken worden over de toepasbaarheid van probabilistische segmentatie op de hyperspectrale beelden van natuurlijke kweldervegetatie op Schiermonnikoog.

7.1. Validatie

De veldgegevens bestaan uit 208 gewone veldopnamen en 176 snelle veldopnamen (zie paragraaf 3.3). Voor training zijn de gewone veldopnamen gebruikt en voor de validatie de snelle veldopnamen. Met behulp van de validatiesamples zijn voor de crisp classificaties confusionmatrices bepaald. Hoe een fuzzy classificatie gevalideerd zou moeten worden is onduidelijk. Voor de hand ligt dat voor goede validatie van een fuzzy classificatie ook fuzzy validatiesamples nodig zouden zijn. Hoe hiermee gevalideerd zou moeten worden is onderwerp voor eventueel aanvullend (literatuur) onderzoek. Daarom zullen de fuzzy classificaties beoordeeld worden op basis van de validatie van de crisp classificaties. Deze crisp classificaties zijn namelijk verkregen door de klasse met de grootste kans of aandeel uit de fuzzy classificatie te selecteren.

Bij het uitgevoerde experiment is probabilistische segmentatie toegepast op de Schiermonnikoog-data met verschillende klassenindelingen. De gebruikte indelingen zijn niveau 1, 2 en 3 en de indeling uit Skidmore e.a. [2001] (zie paragraaf 3.4). Voor de k-nearest-neighbours-classificatie zou volgens de vuistregel (zie subparagraaf 6.3.1) met 208 trainingssamples een k gekozen moeten worden van 14. Omdat twee klassen in niveau 3 echter maar 4 trainingssamples hebben, zou k echter ook niet groter gekozen mogen worden dan 4. Als compromis is daarom $k=7$ gebruikt. Voor de segmentatie zijn de banden 9, 19 en 25 gebruikt, omdat deze ongeveer overeenkomen met de kleuren van de false-colour-luchtfoto's die gebruikt worden voor visuele interpretatie. De gehanteerde drempelwaarden d voor de segmentatiepiramide zijn 2, 3, 4, 5, 6, 7, 8, 9, 10, 12, 14, 16, 20, 24, 28 en 32 (zie paragraaf 5.2.1). De verwaarlozing van kleine klassenaandelen voor het bepalen van het aantal klassen in een segment is vastgesteld op een fractie van het segment: 0,1. Deze verwaarlozing is alleen van invloed op de segmentatie (zie subparagraaf 5.2.3).

Bij het classificeren van één beeld van alle vier de stroken samen (2,5 miljoen pixels, zie figuur 3.3) kunnen maximaal 88 banden gebruikt worden. Bij meer banden wordt één van de tekstbestanden voor het ANN-classificatieprogramma groter dan twee Gigabyte. Dergelijke grote bestanden kunnen bij het gebruikte besturingssysteem niet door alle programma's meer normaal verwerkt worden. Dit had eenvoudig omzeild kunnen worden door de stukken wad van de stroken af te knippen. Aangezien bij gebruik van alle 128 banden waarschijnlijk toch overtraining zou optreden, is het experiment uitgevoerd op 88 banden. Hiervoor zijn de eerste 88 banden gebruikt, waarbij de banden 1, 2, 65, 66 en 67 overgeslagen zijn, omdat deze een nogal slechte signaal-ruisverhouding hebben (zie figuur 7.1). Dit is een nogal arbitraire keuze. Wanneer er verder onderzoek gedaan zou worden naar de mate waarin overtraining zich voordoet verdient het aanbeveling dat hierbij ook nagegaan wordt wat de beste methode voor feature-reductie is (zie paragraaf 6.1).



Figuur 7.1. Spectra van twee willekeurige pixels.

7.2. Resultaten

De validatie van de resultaten laat een verbetering zien door het gebruik van probabilistische segmentatie ten opzichte van crisp k-nearest-neighbours-classificatie (zie onderstaande tabel). Gedetailleerde resultaten zoals figuren en confusionmatrices zijn te vinden in bijlage 1.

Overall-accuracy	Klassenindeling			
	niveau 1 9 klassen	niveau 2 15 klassen	niveau 3 21 klassen	uit Skidmore e.a. [2001] 19 klassen
k-nearest-neighbours-classificatie	61,93 %	46,33 %	40,83 %	48,31 %
probabilistische segmentatie	67,43 %	49,08 %	41,28 %	50,24 %

Ter vergelijking wordt in onderstaande tabel de kwaliteit van enkele andere classificatiemethoden gegeven [Skidmore e.a. 2001 blz. 59].

Overall-accuracy	
false-colour-luchtfoto-interpretatie door operator	43 %
tussenvorm van spectral angle mapper en 1-nearest-neighbour-classificatie	40 %
expertsysteem	66 %

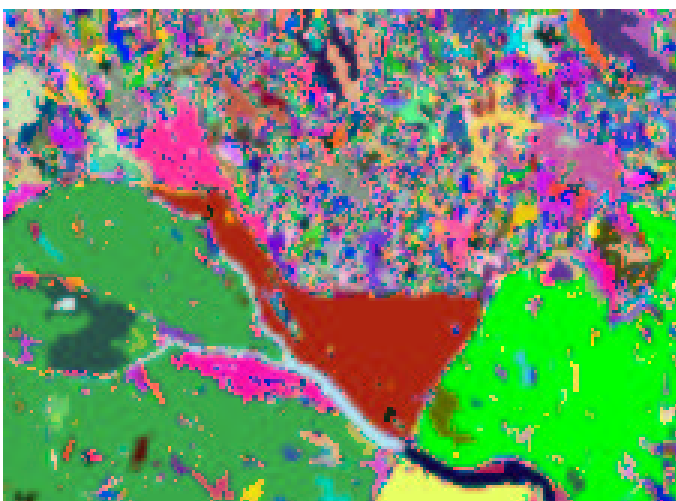
k-Nearest-neighbours-classificatie lijkt een flinke verbetering op te leveren ten opzichte van de 40% overall-accuracy die met de spectral angle mapper gehaald is in Skidmore e.a. [2001]. Er zijn echter enkele verschillen in de gebruikte data die de vergelijking bemoeilijken. Ten eerste zijn in Skidmore e.a. [2001] de geometrisch gecorrigeerde beelden gebruikt in plaats van de ongecorrigeerde. Verder zijn in Skidmore e.a. [2001] de banden niet teruggebracht van 2 naar 1 byte en zijn alle 128 banden gebruikt in plaats van alleen de eerste 88. Een laatste en belangrijk verschil is dat in Skidmore e.a. [2001] voor de klassen water en strand beide slechts 4 samples gebruikt worden. Bij de k-nearest-neighbours-classificatie voor de probabilistische segmentatie zijn voor beide klassen 50 samples gebruikt. Omdat de klassen water en strand goed geassocieerd worden trekt dit de overall-accuracy erg omhoog. Wanneer voor het verschil in het aantal samples voor water en strand gecorrigeerd wordt, komt de overall-accuracy voor k-nearest-neighbours-classificatie uit op 33,54% en de overall-accuracy voor probabilistische segmentatie op 36,02%. Omdat k-nearest-neighbours-classificatie dan slechtere (crisp) resultaten geeft dan de manier waarop in Skidmore e.a. [2001] de spectral angle mapper gebruikt wordt, blijft de overall-accuracy ook na probabilistische segmentatie nog onder de 40%.

7.3. Conclusies en aanbevelingen

Dat de overall-accuracy van k-nearest-neighbours-classificatie onder de 40% blijft, zou veroorzaakt kunnen worden door het gebruik van andere data dan Skidmore e.a. [2001] (minder banden en een lagere radiometrische resolutie). Een andere mogelijke oorzaak is dat k-nearest-neighbours-classificatie, door een grotere gevoeligheid voor overtraining, meer last heeft van het geringe aantal trainingssamples. Onderzocht zou kunnen worden of de tussen-vorm van spectral angle mapper en 1-nearest-neighbour-classificatie die Skidmore e.a. [2001] gebruiken niet alleen betere crisp resultaten oplevert, maar of de kansen ook goed genoeg zijn om te gebruiken voor probabilistische segmentatie.

Hoewel het toepassen van probabilistische segmentatie een verbetering geeft ten opzichte van alleen k-nearest-neighbours-classificatie, is deze verbetering slechts klein. Wat opvalt is dat in sommige gebieden erg veel kleine pure segmenten geselecteerd zijn. Grote delen van de kwelder worden bedekt met kleine segmenten van 1 tot 5 pixels groot, terwijl in andere delen van de kwelder segmenten voorkomen van honderden pixels groot (zie figuur 7.2). Dit was ook de reden om voor een verwaarlozingsfractie van 0,1 te kiezen. Zonder verwaarlozing werden er nog meer kleine segmenten geselecteerd en viel de overall-accuracy van de classificatie ongeveer 1 procent lager uit. Bij het gebruik van één goed gekozen niveau uit de segmentatiepiramide zijn de resultaten bijna even goed als bij selectie van segmenten uit meerdere niveaus. Dit doet vermoeden dat de segmentselectie niet optimaal is. Aan de aanpassingen in de methode ligt het niet want op agrarisch gebied wordt wel een goede segmentatie verkregen. Waar het wel door veroorzaakt wordt zou uit aanvullend onderzoek moeten blijken. Mogelijke oorzaken zijn:

- Te lage kwaliteit van de geschatte kansen met de k-nearest-neighbours-classificatie, veroorzaakt door overtraining (als gevolg van het gebruik van hyperspectrale data bij een gering aantal trainingssamples);
- De aanwezigheid van grote spectrale overlap tussen klassen, al dan niet veroorzaakt door foutieve trainingssamples als gevolg van geometrische onnauwkeurigheid;
- Te lage ruimtelijke resolutie van de beelden ten opzichte van de objectgrootte;
- Een te grote voorkeur van de segmentselectiemethode voor pure segmenten, waardoor een groot fuzzy gebied door kleine pure stukjes daarbinnen wordt opgedeeld in vele kleine segmenten.



Figuur 7.2. Detail van de segmentatie voor klassenindeling niveau 1. Ieder segment heeft een random kleur.

Wanneer de oorzaak te verhelpen is, zal de overall-accuracy van probabilistische segmentatie zeer waarschijnlijk de 40% overtreffen. Zo'n sterke verbetering als het expertsysteem bereikt,

zal niet behaald worden, maar een interessante verbetering zou dan zeker mogelijk moeten zijn. In dat geval zal de geautomatiseerde classificatie door probabilistische segmentatie (zonder het toevoegen van extra informatie) minstens zo nauwkeurig zijn als interpretatie van false-colour-luchtfoto's door een operateur. Het voordeel van probabilistische segmentatie ten opzichte van de huidige productiemethode van vegetatiekaarten zou zijn dat automatische classificatie objectiever is en een fuzzy classificatie oplevert.

De vraag zou gesteld kunnen worden wat de waarde is van een vegetatiekaart met een overall-accuracy van onder de 50%. Voor het detecteren van veranderingen is een classificatie met een dergelijke lage nauwkeurigheid eigenlijk ongeschikt. Zou het niet beter zijn om minder klassen te onderscheiden of de klassen zo te kiezen dat ze spectraal beter te onderscheiden zijn? Hiervoor zouden ook minder uitvoerig veldwerk nodig zijn, wat het mogelijk maakt om meer trainingssamples in te winnen. Hierdoor zou een classificatie verkregen kunnen worden met weliswaar minder ecologische informatie, maar met een nauwkeurigheid die hoog genoeg is om met voldoende zekerheid veranderingen vast te stellen.

Ondanks de matige classificatieresultaten, blijkt probabilistische segmentatie globaal aan de verwachtingen te voldoen door een verbetering van de k-nearest-neighbours-classificatie te geven.

8. Conclusies en aanbevelingen

In dit hoofdstuk zullen de belangrijkste conclusies en aanbevelingen uit het uitgevoerde onderzoek beschreven worden. Hierbij wordt onderscheid gemaakt tussen conclusies en aanbevelingen ten aanzien van fuzzy training voor maximum-likelihood-classificatie (hoofdstuk 4) en het toepassen van probabilistische segmentatie op hyperspectrale beelden van natuurlijke vegetatie.

8.1. Fuzzy training voor maximum-likelihood-classificatie

Conclusies

Met behulp van de vereffeningstheorie en kansmodelschatting is het mogelijk een fuzzy maximum-likelihood-training uit te voeren die uit mixed pixels spectra van pure klassen schat zonder systematische fout. Voordeel van fuzzy training is dat meer pixels in het beeld bruikbaar zijn voor training. Hierdoor is het mogelijk om heterogene gebieden te gebruiken voor training of om de trainingssamples op random plaatsen in te winnen. Voorwaarde voor het gebruik van deze methode voor fuzzy training is dat men beschikt over schattingen van de klassenaandelen van de mixed trainingssamples en dat de klassen in een pixel spectraal lineair vermengd zijn.

Aanbeveling

Om de bruikbaarheid van fuzzy training voor maximum-likelihood-classificatie te beoordelen verdient het aanbeveling om de methode te testen op een beeld waarvoor tijdens veldwerk de klassenaandelen van fuzzy trainingssamples zijn vastgesteld.

8.2. Probabilistische segmentatie

Conclusies

Probabilistische segmentatie is toegepast op de hyperspectrale HyMap-beelden van de natuurlijke kweldervegetatie op Schiermonnikoog. Voor deze beelden en bijbehorende veldgegevens, met een beperkt aantal samples per klasse, geeft de k-nearest-neighbours-classificatie van de ANN-software de best geschatte kansen als invoer voor de probabilistische segmentatie. Deze schattingen zijn beter dan die met maximum-likelihood-classificatie, minimum-distance-classificatie of spectral angle mapper bepaald worden.

Voor toepassing van probabilistische segmentatie op hyperspectrale beelden van natuurlijke vegetatie met een beperkt aantal trainingssamples zijn twee aanpassingen gedaan. Ten eerste is er gekozen voor een implementatie van de k-nearest-neighbours-classificatie die hyperspectrale data kan verwerken. Ten tweede is het segmentselectiecriterium zo aangepast dat de voorkeur niet langer alleen naar pure segmenten uit gaat, maar dat er gestreefd wordt naar segmenten met zo min mogelijk klassen. Hierdoor heeft bijvoorbeeld een segment met twee klassen de voorkeur boven een segment met drie klassen.

Probabilistische segmentatie geeft een lichte verbetering van de overall-accuracy van de classificatie. Deze verbetering is afhankelijk van de klassenindeling. De grootste verbetering treedt op bij een indeling in 9 klassen (niveau 1). De overall-accuracy stijgt hiervoor van 61,93% naar 67,43%. De indeling met de kleinste verbetering heeft 21 klassen (niveau 3). Hiervoor steeg de overall-accuracy van 40,83% naar 41,28%.

Het is onduidelijk waarom de verbetering door probabilistische segmentatie slechts klein is. Dit kan aan de data liggen of aan de probabilistische segmentatiemethode. Een mogelijke oorzaak als gevolg van de data is de spectrale overlap tussen de klassen, maar ook een te lage ruimtelijke resolutie van het beeld of het geringe aantal trainingssamples kan de oorzaak zijn. Dat de segmentselectie van de probabilistische segmentatiemethode een te grote voorkeur heeft voor pure segmenten (ook als deze heel klein zijn) kan ook een oorzaak zijn.

Aanbevelingen

Op de eerste plaats zou onderzocht moeten worden wat de oorzaak is dat de verbetering door probabilistische segmentatie slechts klein is, en hoe het komt dat veel erg kleine segmenten geselecteerd worden. Door het toepassen van probabilistische segmentatie op andere data kan geverifieerd worden of de data het probleem zijn. Door gebruik te maken van een handmatig gedefinieerde segmentatie zou gecontroleerd kunnen worden of de segmentselectiemethode de oorzaak is.

Verder is het wenselijk om te bepalen in welke mate overtraining (zie paragraaf 6.1) zich voordoet bij k-nearest-neighbours-classificatie. In het geval dat overtraining zich voordoet moet een methode gekozen worden om dit te voorkomen.

Ten derde is het sterk aan te raden om een validatiemethode te ontwerpen waarmee een fuzzy classificatie gevalideerd kan worden.

Aanvullend onderzoek naar betere methoden dan k-nearest-neighbours-classificatie om kansen op klassen uit hyperspectrale data te schatten zou wenselijk zijn. Eén van de te onderzoeken methode hiervoor is het bepalen van kansen uit de spectrale hoek naar het dichtstbijzijnde trainingssample zoals voor het expertsysteem [Skidmore e.a. 2001] gedaan is.

Tenslotte zou overwogen kunnen worden om het expertsysteem en probabilistische segmentatie te integreren door de kansen die het expertsysteem bepaald te gebruiken voor probabilistische segmentatie.

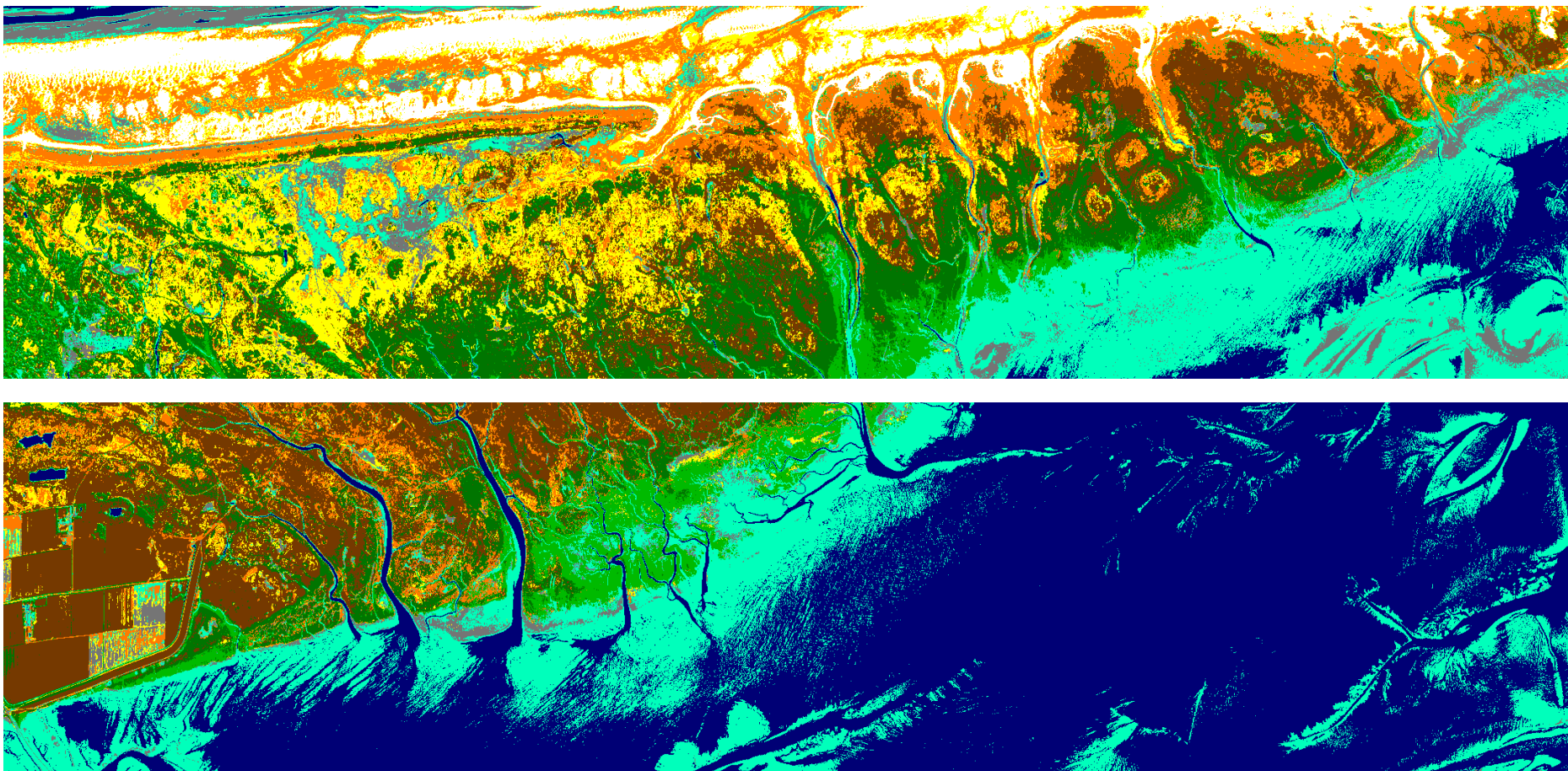
Literatuur

- [Clark 1999] R.N. Clark, "Spectroscopy of Rocks and Minerals, and Principles of Spectroscopy". In: A. Rencz (redacteur), *Manual of Remote Sensing*. New York: John Wiley and Sons, 1999.
- [Droesen 1999] W. Droesen, *Spatial modelling and monitoring of natural landscapes; With cases in the Amsterdam Waterworks Dunes*. Proefschrift Landbouwniversiteit Wageningen, 1999.
- [Duin 2001] R.P.W. Duin, *Handouts Pattern recognition course tn3534*. Technische Universiteit Delft, Afdeling Technische Natuurkunde, 2001.
ftp://ftp.ph.tn.tudelft.nl/pub/bob/PR_Course/handouts/sheets_19nov01_color.pdf
Bezocht november 2002.
- [Eastman en Laney 2002] J.R. Eastman en R.M. Laney. "Bayesian Soft Classification for Sub-Pixel Analysis: A Critical Evaluation". *Photogrammetric Engineering & Remote Sensing*, jaargang 68, nr. 11 (2002), blz. 1149-1154.
- [ENVI Tutorials 1999] *ENVI Tutorials; version 3.2*. Better Solutions Consulting, 1999.
- [Fisher 1997] P. Fisher, "The pixel: a snare and a delusion". *International Journal of Remote Sensing*, jaargang 18, nr. 3 (1997), blz. 679-685.
- [Foody 1999] G.M. Foody. "The continuum of classification fuzziness in thematic mapping". *Photogrammetric Engineering & Remote Sensing*, jaargang 65, nr. 4 (1999), blz. 443-451.
- [Gorte 1998] B.G.H. Gorte, *Probabilistic segmentation of remotely sensed images*. Proefschrift Landbouwniversiteit Wageningen en ITC, 1998.
- [Hoekstra 1998] A. Hoekstra, *Generalisation in Feed Forward Neural Classifiers*. Proefschrift Technische Universiteit Delft, 1998.
- [Janssen 2001] J.A.M. Janssen, *Monitoring of salt-marsh vegetation by sequential mapping*. Proefschrift Universiteit van Amsterdam en Meetkundige Dienst van Rijkswaterstaat, 2001.
- [Mount 1998] D.M. Mount, *ANN Programming Manual*. University of Maryland, 1998.
- [Skidmore e.a. 2001] A.K. Skidmore, K. Schmidt, H. Kloosterman, L. Kumar en H. van Oosten, *Hyperspectral imagery for coastal wetland vegetation mapping*. BeleidsCommissie Remote Sensing, Meetkundige Dienst van Rijkswaterstaat, 2001.
- [Teunissen 1994] P.J.G. Teunissen, *Mathematische geodesie I (ge30); Inleiding vereffenings-theorie*. Collegedictaat Technische Universiteit Delft, Faculteit der Geodesie, 1994.
- [Verhoef 1997] H.M.E. Verhoef, *Geodetische Deformatie Analyse (ge135)*. Collegedictaat Technische Universiteit Delft, Faculteit der Geodesie, 1997.
- [Wang 1990] F. Wang, "Fuzzy supervised classification of remote sensing images". *IEEE Transactions on Geoscience and Remote Sensing*, jaargang 28, nr. 2 (1990), blz. 194-201.

Bijlage 1. Classificatieresultaten

In deze bijlage zijn enkele resultaten opgenomen van de toepassing van probabilistische segmentatie op de Schiermonnikoog-data (zie hoofdstuk 2). Hiervoor zijn de onderstaande parameters gebruikt.

- Klassenindelingen: niveau 1, 2 en 3 en Skidmore (zie paragraaf 3.4).
- k-Nearest-neighbours-classificatie: $k=7$.
- Drempelwaarden voor de segmentatiepiramide: 2, 3, 4, 5, 6, 7, 8, 9, 10, 12, 14, 16, 20, 24, 28 en 32.
- De verwaarlozing van kleine klassenaandelen voor het bepalen van het aantal klassen in een segment: fractie 0,1 (zie paragraaf 5.2.3).
- 88 banden: nummer 3 tot en met 64 en nummer 68 tot en met 93.



Figuur B1.1. Crisp k-nearest-neighbours-classificatie van stroken 2 en 4 voor klassenindeling niveau 1. Opvallend is de overgang tussen strook 2 en 4 als gevolg van radiometrische verschillen tussen de stroken.

Klassenindeling niveau 1

Crisp k-nearest-neighbours-classificatie

(zie figuur B1.1)

Zelfs bij dit beperkt aantal klassen treden nog relatief veel classificatiefouten op. Er vindt flinke verwarring plaats tussen de klassen middelhoge (4) en hoge kwelder (5); Wat "middelhoge kwelder" is, wordt vaak geclassificeerd als hoge kwelder. Duin (7) wordt vaak als "overgang duin kwelder" (6) geclassificeerd. Terwijl de klasse "kale grond en ijle vegetatie" (1) bijna altijd als pionierkwelder (2) geclassificeerd wordt. Opvallend is dat de klassen waar veel samples voor zijn beter geclassificeerd worden.

Confusionmatrix

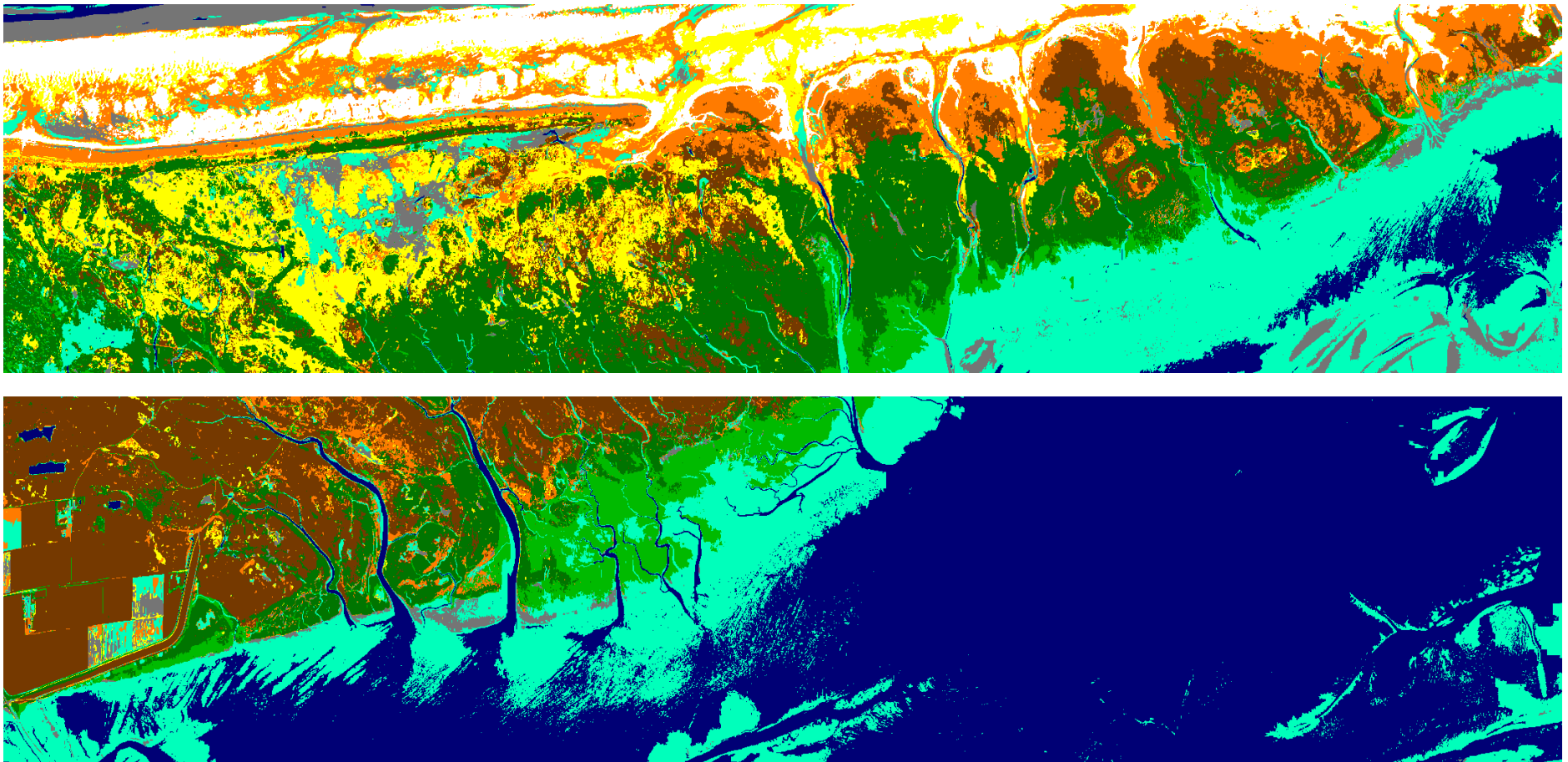
	1	2	3	4	5	6	7	8	9	
1	1	4	1	0	0	0	0	0	0	0,17
2	6	5	0	1	1	1	1	0	0	0,33
3	0	2	4	6	0	3	0	0	0	0,27
4	0	0	4	35	1	2	0	0	0	0,83
5	1	0	0	17	20	0	5	0	0	0,47
6	0	2	5	6	0	18	6	0	0	0,49
7	0	1	0	4	1	1	2	0	0	0,22
8	0	0	0	0	0	1	0	25	0	0,96
9	0	0	0	0	0	0	0	0	25	1,00
	0,12	0,36	0,29	0,51	0,87	0,69	0,14	1,00	1,00	

gemiddelde accuracy: 55,33%

overall-accuracy: 61,93%

Legenda voor crisp classificaties van klassenindeling niveau 1

- 1  Kale grond en ijle vegetatie
- 2  Pionierkwelder
- 3  Lage kwelder
- 4  Middelhoge kwelder
- 5  Hoge kwelder
- 6  Overgang duin kwelder
- 7  Duin
- 8  Strand
- 9  Water



Figuur B1.2. Crisp classificatie per segment van stroken 2 en 4 voor klassenindeling niveau 1. Een duidelijk verschil met figuur B1.1 is dat deze classificatie een minder ruizige kaart oplevert.

Klassenindeling niveau 1

Crisp classificatie per segment

(zie figuur B1.2).

De klassen middelhoge kwelder (4) en "overgang duin kwelder" (6) worden beter geclassificeerd dan bij crisp k-nearest-neighbours-classificatie. Hoge kwelder (5) heeft daarentegen juist een hogere classificatiefout. De overige klassen worden even goed geclassificeerd, of net iets beter (één sample) geclassificeerd.

Confusionmatrix

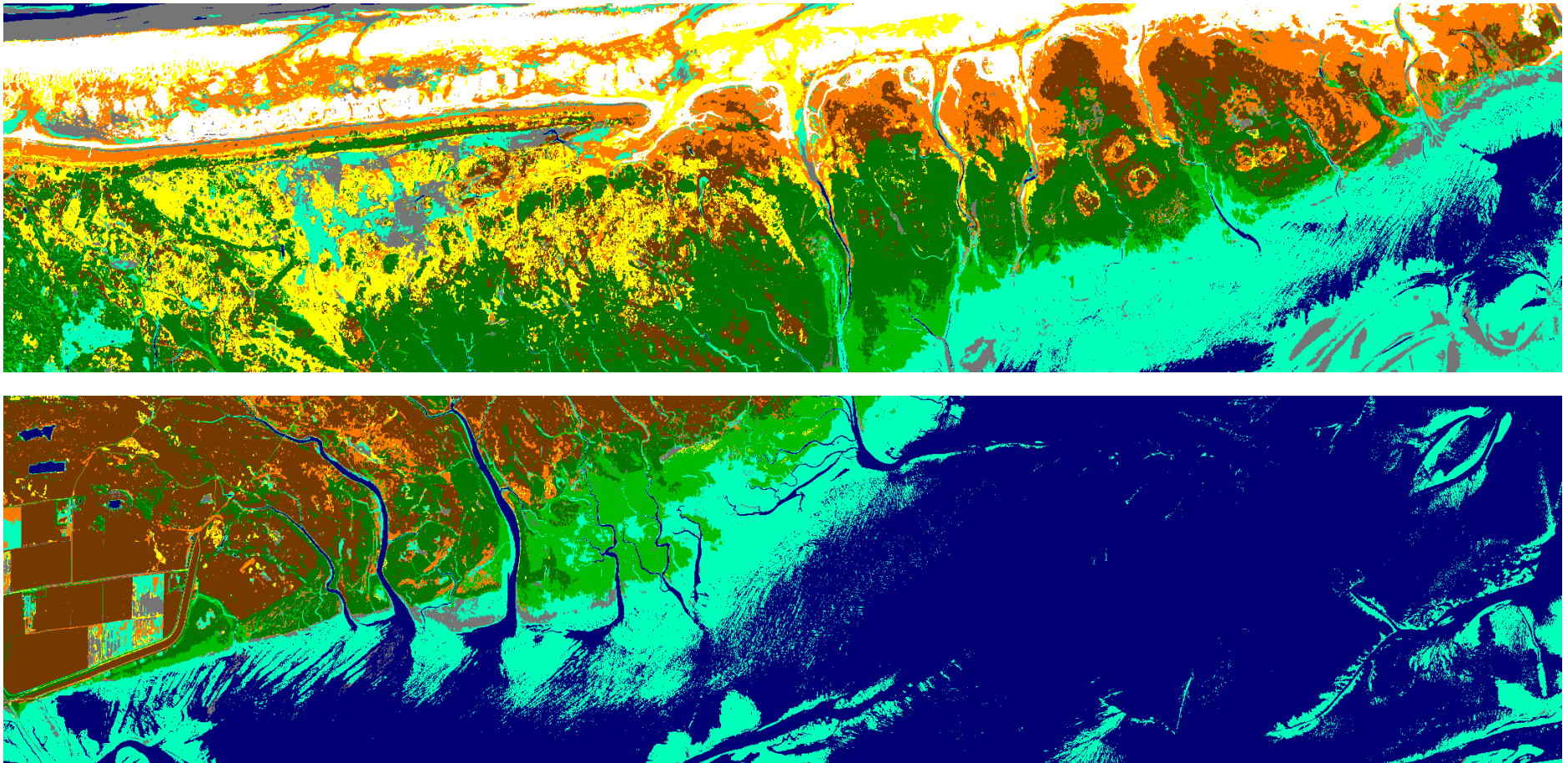
	1	2	3	4	5	6	7	8	9	
1	2	2	0	0	0	0	0	0	0	0,50
2	5	5	0	1	1	1	0	0	0	0,38
3	0	4	4	1	0	0	0	0	0	0,44
4	0	1	5	47	8	2	0	0	0	0,75
5	1	0	0	12	12	0	5	0	0	0,40
6	0	1	5	4	0	21	6	0	0	0,57
7	0	0	0	4	2	1	3	0	0	0,30
8	0	1	0	0	0	1	0	25	0	0,93
9	0	0	0	0	0	0	0	0	25	1,00
	0,25	0,36	0,29	0,68	0,52	0,81	0,21	1,00	1,00	

gemiddelde accuracy: 56,86%

overall-accuracy: 66,06%

Legenda voor crisp classificaties van klassenindeling niveau 1

- 1  Kale grond en ijle vegetatie
- 2  Pionierkwelder
- 3  Lage kwelder
- 4  Middelhoge kwelder
- 5  Hoge kwelder
- 6  Overgang duin kwelder
- 7  Duin
- 8  Strand
- 9  Water



Figuur B1.3. Crisp classificatie met probabilistische segmentatie van stroken 2 en 4 voor klassenindeling niveau 1. Ten opzichte van figuur B1.2 komt in sommige gebieden iets van de ruis uit figuur B1.1 in deze classificatie terug. In andere gebieden blijft deze afwezig.

Klassenindeling niveau 1

Crisp classificatie met probabilistische segmentatie

(zie figuur B1.3)

De klassen middelhoge kwelder (4) en "overgang duin kwelder" (6) worden beter geclassificeerd dan bij crisp k-nearest-neighbours-classificatie. Hoge kwelder (5) heeft daarentegen juist een hogere classificatiefout. De overige klassen worden even goed geclassificeerd of net iets beter (één sample) geclassificeerd.

Ten opzichte van de crisp classificatie per segment wordt de klasse middelhoge kwelder (4) nog beter geclassificeerd, terwijl de klasse "overgang duin kwelder" (6) minder vaak slechter geclassificeerd wordt

Confusionmatrix

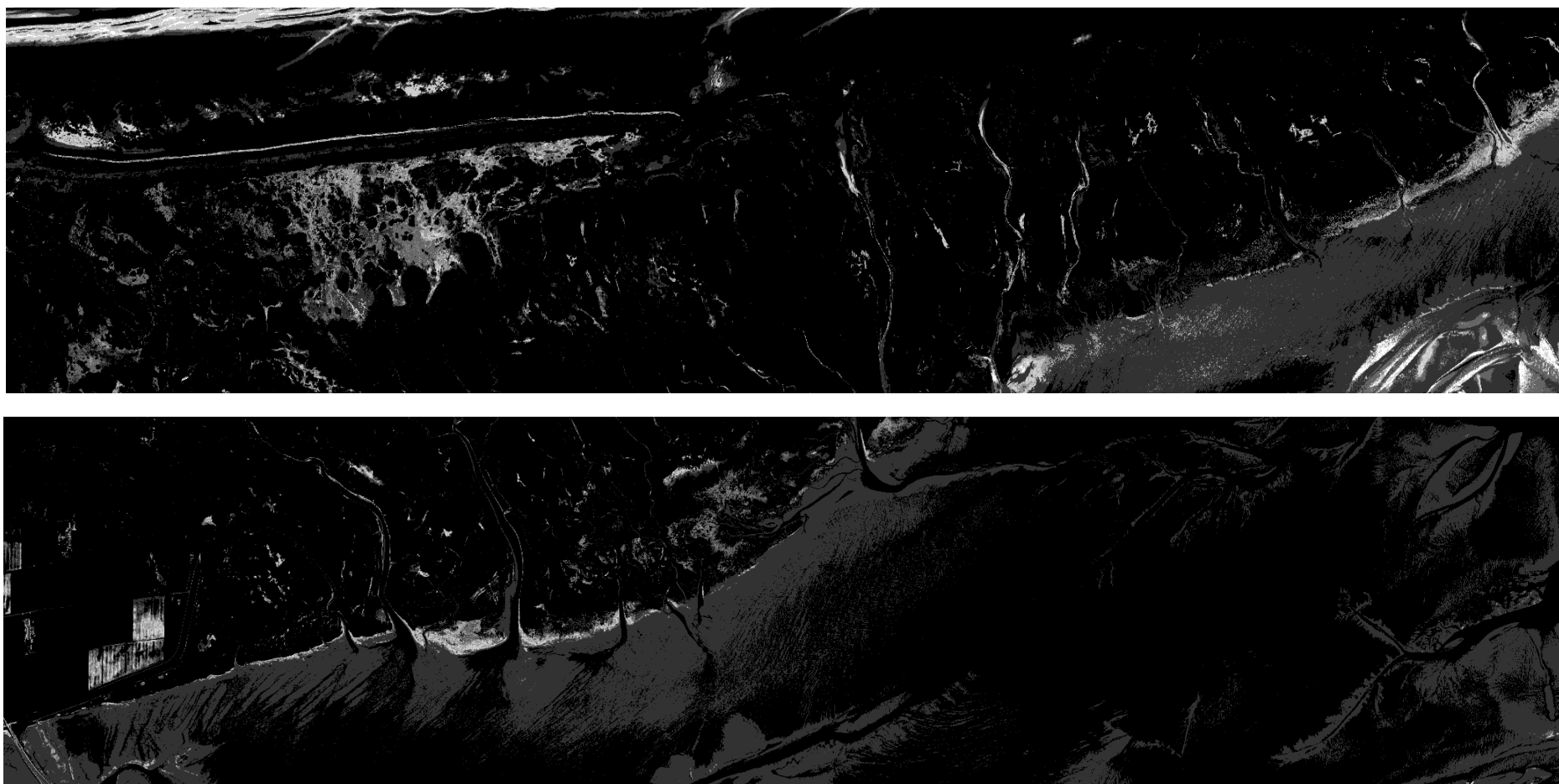
	1	2	3	4	5	6	7	8	9	
1	2	2	0	0	1	0	1	0	0	0,33
2	5	5	0	1	1	1	0	0	0	0,38
3	0	4	4	1	0	0	0	0	0	0,44
4	0	1	5	48	6	2	0	0	0	0,77
5	1	0	0	12	14	0	5	0	0	0,44
6	0	1	5	4	0	21	5	0	0	0,58
7	0	0	0	3	1	1	3	0	0	0,38
8	0	1	0	0	0	1	0	25	0	0,93
9	0	0	0	0	0	0	0	0	25	1,00
	0,25	0,36	0,29	0,70	0,61	0,81	0,21	1,00	1,00	

gemiddelde accuracy: 57,99%

overall-accuracy: 67,43%

Legenda voor crisp classificaties van klassenindeling niveau 1

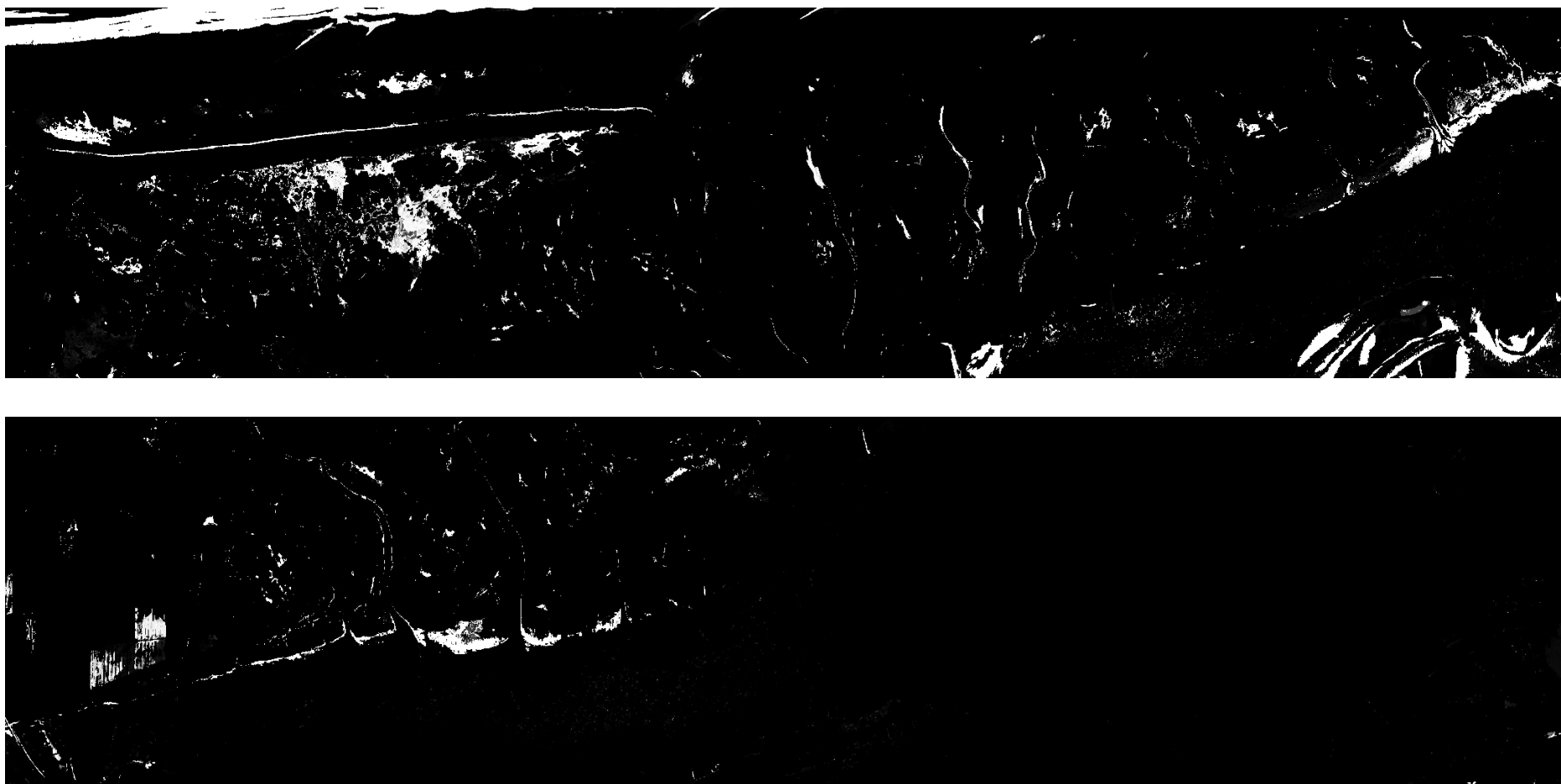
- 1  Kale grond en ijle vegetatie
- 2  Pionierkwelder
- 3  Lage kwelder
- 4  Middelhoge kwelder
- 5  Hoge kwelder
- 6  Overgang duin kwelder
- 7  Duin
- 8  Strand
- 9  Water



Figuur B1.4. Fuzzy k-nearest-neighbours-classificatie van stroken 2 en 4 voor klasse de "kale grond en ijle vegetatie" (1) uit klassenindeling niveau 1. Opvallend is dat de kans op deze klasse niet boven de 0,45 uit komt door de spectrale verwarring met andere klassen.

Legenda

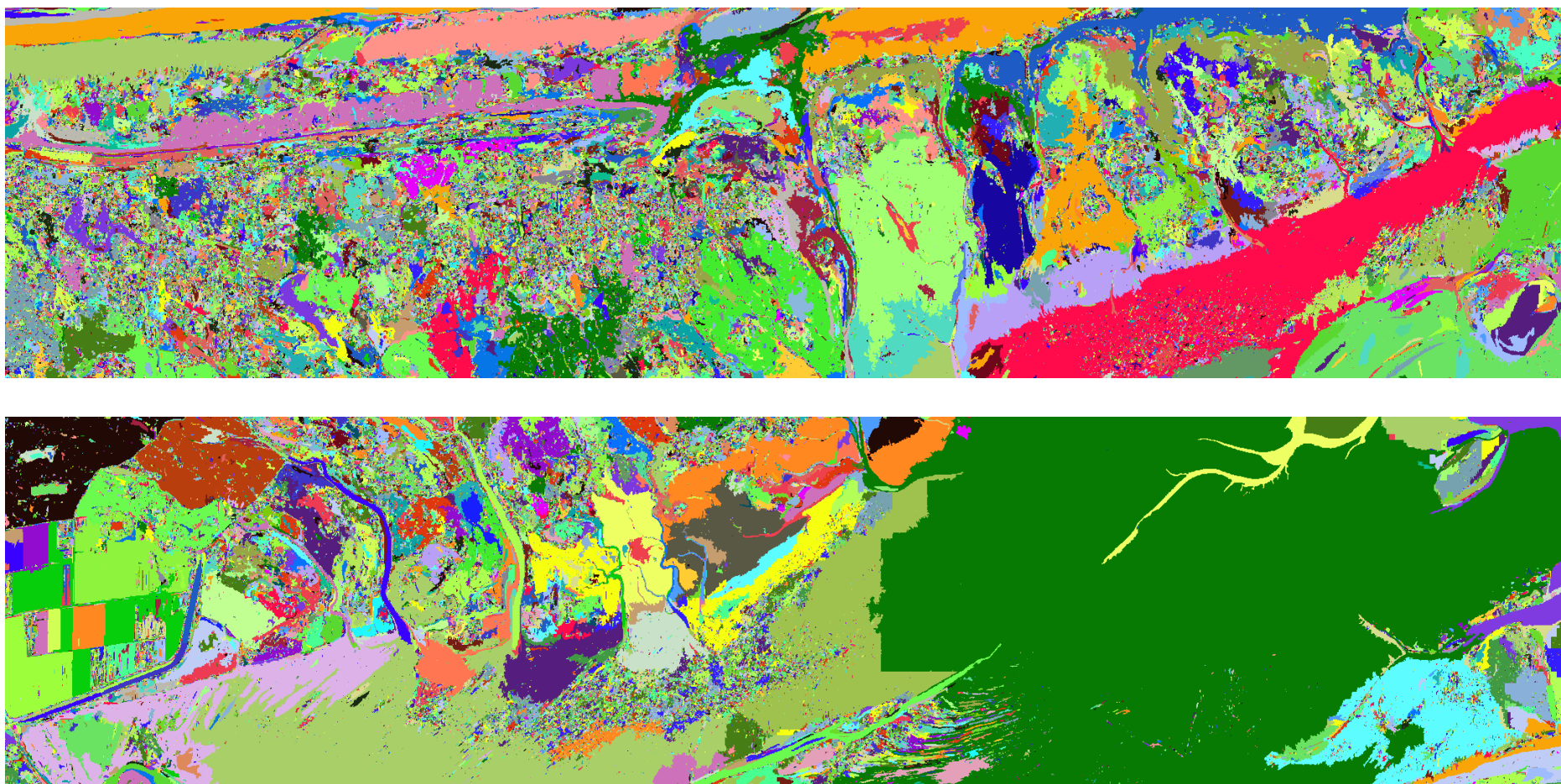




Figuur B1.5. Fuzzy classificatie met probabilistische segmentatie van stroken 2 en 4 voor klasse de "kale grond en ijle vegetatie" (1) uit klassenindeling niveau 1. Omdat voor ieder segment enkele klassen uitgesloten worden, zijn de geschatte kansen meer uitgesproken dan in figuur B1.4. Overigens zijn niet alle kansen exact gelijk aan 0 of 1.

Legenda





Figuur B1.6. Segmentatiebestand van stroken 2 en 4 voor klassenindeling niveau 1. Ieder segment heeft een random kleur. In sommige gebieden worden veel erg kleine segmenten gekozen.

Klassenindeling niveau 2

crisp k-nearest-neighbours-classificatie

Confusionmatrix

	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	
1 Kale grond en ijle vegetatie	1	4	0	0	0	0	0	0	0	0	0	0	0	0	0	0.20
2 Salicornia-Suaeda-vegetatie	6	3	0	0	0	1	0	0	0	1	0	1	1	0	0	0.23
3 Puccinellia-vegetatie	0	0	1	0	0	1	0	0	0	0	0	5	0	0	0	0.14
4 Limonium-vegetatie	0	2	0	0	0	2	5	1	0	0	0	0	0	0	0	0.00
5 Halimione-vegetatie	0	3	2	0	1	0	0	0	1	0	0	0	0	0	0	0.14
6 Juncus-ger.-vegetatie	0	0	0	0	0	4	1	0	0	1	0	0	0	0	0	0.67
7 Festuca-vegetatie	0	0	2	0	0	2	7	0	0	0	0	0	0	0	0	0.64
8 Juncus-mar.-vegetatie	0	0	2	0	0	5	0	5	8	1	1	0	0	0	0	0.23
9 Elymus-vegetatie	0	0	1	0	0	3	1	2	5	3	1	4	1	0	0	0.24
10 Potentilla-vegetatie	0	0	0	0	0	6	0	0	0	5	1	0	1	0	0	0.38
11 Sagina-vegetatie	1	0	2	0	0	1	1	0	1	5	3	0	3	0	0	0.18
12 Overgang duin kwelderveg.	0	1	3	0	0	1	0	0	2	0	0	14	6	0	0	0.52
13 Duin	0	1	0	0	0	1	0	1	1	1	0	1	2	0	0	0.25
14 Strand	0	0	0	0	0	0	0	0	0	0	0	1	0	25	0	0.96
15 Water	0	0	0	0	0	0	0	0	0	0	0	0	0	0	25	1.00
	0.12	0.21	0.08	?	1.00	0.15	0.47	0.56	0.28	0.29	0.50	0.54	0.14	1.00	1.00	

gemiddelde accuracy: 45,28%

overall-accuracy: 46,33%

Crisp classificatie met probabilistische segmentatie

Confusionmatrix

	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	
1 Kale grond en ijle vegetatie	2	2	0	0	0	0	0	0	0	0	0	0	1	0	0	0.40
2 Salicornia-Suaeda-vegetatie	5	5	0	0	0	0	0	0	0	1	0	1	0	0	0	0.42
3 Puccinellia-vegetatie	0	0	1	0	0	2	0	0	0	0	0	2	0	0	0	0.20
4 Limonium-vegetatie	0	2	0	0	0	2	3	1	0	0	0	0	0	0	0	0.00
5 Halimione-vegetatie	0	3	2	0	1	0	0	0	1	0	0	0	0	0	0	0.14
6 Juncus-ger.-vegetatie	0	0	0	0	0	6	1	0	1	1	0	0	0	0	0	0.67
7 Festuca-vegetatie	0	0	2	0	0	3	9	0	0	1	0	0	0	0	0	0.60
8 Juncus-mar.-vegetatie	0	0	2	0	0	4	0	3	6	1	0	0	0	0	0	0.19
9 Elymus-vegetatie	0	0	2	0	0	3	1	4	5	5	2	4	0	0	0	0.19
10 Potentilla-vegetatie	0	0	0	0	0	4	0	0	0	3	2	0	1	0	0	0.30
11 Sagina-vegetatie	1	0	0	0	0	1	1	0	2	5	2	0	4	0	0	0.12
12 Overgang duin kwelderveg.	0	1	4	0	0	1	0	0	2	0	0	17	5	0	0	0.57
13 Duin	0	0	0	0	0	1	0	1	1	0	0	1	3	0	0	0.43
14 Strand	0	1	0	0	0	0	0	0	0	0	0	1	0	25	0	0.93
15 Water	0	0	0	0	0	0	0	0	0	0	0	0	0	0	25	1.00
	0.25	0.36	0.08	?	1.00	0.22	0.60	0.33	0.28	0.18	0.33	0.65	0.21	1.00	1.00	

gemiddelde accuracy: 46,40%

overall-accuracy: 49,08%

Klassenindeling niveau 3

Crisp k-nearest-neighbours-classificatie

Confusionmatrix

	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	16	17	18	19	20	21	
1	1	2	1	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0,25
2	4	0	2	0	0	0	0	0	0	0	0	0	0	0	1	0	0	0	0	0	0	0,00
3	2	0	2	0	0	1	0	0	2	0	0	0	0	0	0	0	0	5	2	0	0	0,14
4	0	1	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0,00
5	0	0	0	1	1	0	0	0	2	0	0	0	0	1	0	0	0	0	0	0	0	0,20
6	0	0	0	0	1	0	0	0	1	0	0	0	0	0	0	0	0	8	2	0	0	0,00
7	0	0	0	1	0	0	0	0	0	1	0	4	1	0	0	0	0	0	0	0	0	0,00
8	0	1	0	1	2	0	0	1	0	0	0	0	0	0	0	0	0	0	0	0	0	0,20
9	0	0	0	0	0	0	0	0	1	2	0	4	0	0	0	0	0	0	0	0	0	0,14
10	0	0	0	0	0	0	0	0	2	2	0	0	0	0	1	0	0	0	0	0	0	0,40
11	1	0	0	0	1	1	0	0	1	2	1	3	0	1	0	1	1	0	0	0	0	0,08
12	0	0	0	0	1	0	0	0	0	0	0	2	0	0	0	0	0	0	0	0	0	0,67
13	0	0	0	0	1	1	0	0	1	2	0	0	5	7	0	1	0	0	0	0	0	0,28
14	0	0	0	0	0	1	0	0	1	0	0	1	2	5	1	1	1	4	0	0	0	0,29
15	0	0	0	0	0	0	0	0	0	5	0	0	0	0	0	1	0	0	1	0	0	0,00
16	0	0	0	0	0	0	0	0	0	1	0	0	0	0	1	6	1	0	1	0	0	0,60
17	0	0	0	0	0	0	0	0	0	0	0	0	0	1	0	3	3	0	3	0	0	0,30
18	0	0	0	1	1	1	0	0	0	0	0	0	0	2	0	0	0	7	3	0	0	0,47
19	0	1	0	0	0	0	0	0	1	0	0	0	1	1	0	0	0	1	2	0	0	0,29
20	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	1	0	25	0	0,96
21	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	25	1,00
	0,12	0,00	0,40	0,00	0,12	0,00	?	1,00	0,08	0,13	1,00	0,14	0,56	0,28	0,00	0,46	0,50	0,27	0,14	1,00	1,00	

gemiddelde accuracy: 36,08%

overall-accuracy: 40,83%

Crisp classificatie met probabilistische segmentatie

Confusionmatrix

	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	16	17	18	19	20	21	
1	1	2	1	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0,25
2	5	0	1	0	0	0	0	0	0	0	0	0	0	0	1	0	0	0	0	0	0	0,00
3	1	0	2	0	0	1	0	0	1	0	0	0	0	0	0	0	0	5	1	0	0	0,18
4	0	1	1	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0,00
5	0	0	0	2	2	0	0	0	0	0	0	0	0	1	0	0	0	0	0	0	0	0,40
6	0	0	0	0	1	0	0	0	1	0	0	0	0	0	0	0	0	7	1	0	0	0,00
7	0	0	0	0	0	0	0	0	1	1	0	3	1	0	0	0	0	0	0	0	0	0,00
8	0	1	0	0	1	0	0	1	0	0	0	0	0	0	0	0	0	0	0	0	0	0,33
9	0	0	0	0	0	0	0	0	1	5	0	4	0	0	0	0	0	0	0	0	0	0,10
10	0	0	0	0	0	0	0	0	0	2	0	0	0	0	1	0	0	0	0	0	0	0,67
11	1	0	0	0	1	1	0	0	1	1	1	1	0	0	0	0	0	0	0	0	0	0,14
12	0	0	0	1	1	0	0	0	0	0	0	5	0	0	0	0	0	0	0	0	0	0,71
13	0	0	0	0	1	1	0	0	2	2	0	0	3	7	0	1	0	0	0	0	0	0,18
14	0	0	0	0	0	1	0	0	3	0	0	1	4	5	1	4	2	4	0	0	0	0,20
15	0	0	0	0	0	0	0	0	0	3	0	0	0	0	0	0	0	0	0	0	0	0,00
16	0	0	0	0	0	0	0	0	0	1	0	0	0	1	1	4	2	0	1	0	0	0,40
17	0	0	0	0	0	0	0	0	0	0	0	0	0	1	0	4	2	0	4	0	0	0,18
18	0	0	0	1	1	1	0	0	1	0	0	0	0	0	0	0	0	8	4	0	0	0,50
19	0	0	0	0	0	0	0	0	1	0	0	0	1	3	0	0	0	1	3	0	0	0,33
20	0	1	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	1	0	25	0	0,93
21	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	25	1,00
	0,12	0,00	0,40	0,00	0,25	0,00	?	1,00	0,08	0,13	1,00	0,36	0,33	0,28	0,00	0,31	0,33	0,31	0,21	1,00	1,00	

gemiddelde accuracy: 35,61%

overall-accuracy: 41,28%

Klassenindelingen met verschillende mate van aggregatie			
Niveau 1	Niveau 2	Niveau 3	Subtypen
1 Kale grond en ijle veg.	1 Kale grond en ijle veg.	1 Kale grond en ijle vegetatie	O + 1.1 + 2.1
2 Pionierkwelder	2 Salicornia-Suaeda-veg.	2 Salicornia (high cover)	2.2
		3 Salicornia-Suaeda	3.1
		4 Suaeda-Limonium	3.2
3 Lage kwelder	3 Puccinellia-vegetatie	5 Puccinellia-Salicornia-Limonium	4.1 + 4.2
		6 Puccinellia-Glaux	4.3 + 4.4
	4 Limonium-vegetatie	7 Limonium	5.1
	5 Halimione-vegetatie	8 Halimione	6.1 + 6.2
4 Middelhoge kwelder	6 Juncus-ger.-vegetatie	9 Juncus-Suaeda-Limonium	8.1 + 8.2
		10 Juncus-Festuca	8.3 + 8.4
	7 Festuca-vegetatie	11 Festuca	9.1 + 9.2
		12 Artemisia	9.3abc
	8 Juncus-mar.-vegetatie	13 Juncus mar.	10.1 + 10.2
	9 Elymus-vegetatie	14 Elymus	11.1 + 11.2
5 Hoge kwelder	10 Potentilla-vegetatie	15 Potentilla-Juncus	12.1
		16 Agrostis-Potentilla	12.2 + 12.3 + 12.4
	11 Sagina-vegetatie	17 Sagina-Armeria	13.1
6 Overgang duin kwelder	12 Overgang duin kwelder	18 Overgang duin kwelder	14.1 + 15.1 + 15.2
7 Duin	13 Duin	19 Duin	D
8 Strand	14 Strand	20 Strand	B
9 Water	15 Water	21 Water	W

Klassenindeling zoals gebruikt in Skidmore e.a. [2001]

Crisp k-nearest-neighbours-classificatie

Confusionmatrix

	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	16	17	18	19	
1 Water	25	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	1.00
2 Bare(SpaSal)	0	2	4	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0.33
3 SuSa	0	4	3	0	0	0	0	0	0	0	0	0	0	0	0	0	0	2	0	0.33
4 SaLim	0	0	2	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0.00
5 Lim	0	0	0	0	0	0	0	4	0	1	0	2	0	0	0	0	0	0	0	0.00
6 Pucc	0	0	1	0	0	1	0	0	0	0	0	0	0	0	0	0	4	0	0	0.17
7 Hal	0	0	1	0	0	1	1	0	1	0	0	0	0	0	0	0	0	0	0	0.25
8 FestArte	0	0	0	0	0	1	0	5	0	1	0	0	0	0	0	0	0	0	0	0.71
9 JunSuaLim	0	0	0	0	0	0	0	3	2	7	0	0	0	1	0	0	0	0	0	0.15
10 JunAtrFes	0	0	0	0	0	0	0	0	2	2	0	1	1	1	0	0	0	0	0	0.29
11 Fest	0	0	0	0	0	3	0	1	0	0	0	0	2	0	0	0	0	0	0	0.00
12 JunM	0	0	0	0	0	1	0	0	1	2	0	3	0	0	0	0	0	0	0	0.43
13 Ely	0	0	0	0	0	2	0	1	2	0	0	2	4	1	2	1	4	0	0	0.21
14 Agros	0	0	0	0	0	0	0	0	0	0	0	0	0	4	0	2	0	2	0	0.50
15 AgrosLot	0	0	0	0	0	0	0	0	0	1	0	1	2	0	2	0	0	2	0	0.25
16 SagArm	0	0	0	0	0	1	0	0	0	0	1	0	5	4	0	3	0	3	0	0.18
17 Dunefoot	0	0	0	1	0	3	0	0	0	0	0	0	1	0	0	0	17	4	0	0.65
18 Dune	0	0	0	0	0	0	0	0	0	0	0	0	2	0	0	0	0	1	0	0.33
19 Beach	0	0	1	0	0	0	0	0	0	0	0	0	0	0	0	0	1	0	25	0.93
	1.00	0.33	0.25	0.00	?	0.08	1.00	0.36	0.25	0.14	0.00	0.33	0.24	0.36	0.50	0.50	0.65	0.07	1.00	

gemiddelde accuracy: 39,27%

overall-accuracy: 48,31%

Crisp classificatie met probabilistische segmentatie

Confusionmatrix

	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	16	17	18	19	
1 Water	25	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	1.00
2 Bare(SpaSal)	0	2	4	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0.33
3 SuSa	0	4	2	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0.33
4 SaLim	0	0	1	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0.00
5 Lim	0	0	0	0	0	0	0	3	0	1	0	0	0	0	0	0	0	0	0	0.00
6 Pucc	0	0	1	0	0	1	0	0	0	0	0	0	0	0	0	0	1	0	0	0.33
7 Hal	0	0	2	0	0	3	1	0	1	0	0	0	0	0	0	0	0	0	0	0.14
8 FestArte	0	0	0	0	0	1	0	6	0	0	0	0	0	0	0	0	0	0	0	0.86
9 JunSuaLim	0	0	0	0	0	0	0	3	3	7	0	1	0	0	0	0	0	0	0	0.21
10 JunAtrFes	0	0	0	0	0	0	0	0	1	2	0	0	1	1	0	0	0	0	0	0.40
11 Fest	0	0	0	0	0	2	0	1	0	0	0	0	1	0	0	0	0	0	0	0.00
12 JunM	0	0	0	0	0	0	0	0	1	2	0	4	2	0	0	0	0	0	0	0.44
13 Ely	0	0	0	0	0	2	0	1	2	0	0	3	4	1	2	2	4	0	0	0.19
14 Agros	0	0	0	0	0	0	0	0	0	2	0	0	0	3	0	2	0	1	0	0.38
15 AgrosLot	0	0	0	0	0	0	0	0	0	0	0	0	2	0	2	0	0	0	0	0.50
16 SagArm	0	0	0	0	0	1	0	0	0	0	1	0	4	6	0	2	0	4	0	0.11
17 Dunefoot	0	0	1	1	0	3	0	0	0	0	0	0	1	0	0	0	19	6	0	0.61
18 Dune	0	0	0	0	0	0	0	0	0	0	0	1	2	0	0	0	1	3	0	0.43
19 Beach	0	0	1	0	0	0	0	0	0	0	0	0	0	0	0	0	1	0	25	0.93
	1.00	0.33	0.17	0.00	?	0.08	1.00	0.43	0.38	0.14	0.00	0.44	0.24	0.27	0.50	0.33	0.73	0.21	1.00	

gemiddelde accuracy: 40,30%

overall-accuracy: 50,24%

Bijlage 2. Handleiding voor de programmatuur

Voor het onderzoek is gebruik gemaakt van de probabilistische-segmentatie-software die ontwikkeld is door B. Gorte, bestaande uit C-programma's en bash-scripts. Veel onderdelen zijn echter voor dit onderzoek aangepast. Verder zijn aanvullende programma's geschreven en zijn afzonderlijke programma's samen in een script opgenomen. Daarnaast is voor dit onderzoek nog gebruik gemaakt van een classificatieprogramma (ANN) dat ontwikkeld is bij de Universiteit van Maryland. Tenslotte worden nog enkele standaard unix-programma's gebruikt zoals paste en xv. In deze bijlage zal het gebruik beschreven worden van de programmatuur die nodig is om de resultaten te reproduceren. De cursief gedrukte commando's worden toegelicht in de gelijknamige paragraaf.

Rechten

Probabilistische segmentatie (probseg)

Auteurs: B.G.H. Gorte
J. Lesparre
Copyright: ITC, Enschede
MD, Delft
TU-Delft, Delft.

Approximate nearest neighbours (ANN)

Auteurs: D.M. Mount
S. Arya
Copyright: Universiteit van Maryland, College Park, VS

Installeren

De C-programma's, bash-scripts en test data zijn te vinden in de map /probseg/* . Van de C-programma's in de map /ann/* is alleen ann_sample benodigd. Kopieer alle bestanden naar een gewenste locatie. Compileer indien nodig alle C-programma's (*.c), bijvoorbeeld met gcc. Vergeet niet het recht op uitvoeren in te stellen met het commando chmod, en indien gewenst de map met de programma's op te nemen in de .profile.

Gebruik

De software gebruikt het dataformaat van ILWIS 1.41 voor beelden en tabellen. In dit formaat bestaat een beeld uit twee bestanden per band: een databestand met één of meer bytes per pixel (*.mpd) en een header-bestand met informatie zoals het aantal rijen en kolommen (*.mpi). Het aantal rijen en kolommen kan opgevraagd worden met het commando *mpi* [naam] nl nc. De mpd-bestanden kunnen vervaardigd worden door de meeste beeldverwerkingsprogramma's of verkregen worden door een band-sequential-bestand (*.bsq in Erdas Imagine) in stukken te knippen. De mpi-bestanden kunnen gemaakt worden met het programma *makempi* [naam] [#kol] [#rij] [#byte]. Omdat deze mpi-bestanden voor alle banden gelijk zijn, kan er één met *makempi* gemaakt worden en gekopieerd worden voor de andere banden.

De software kan alleen bestanden met één byte per pixel aan. Een uitzondering hierop zijn segmentatiebestanden. Wanneer de te gebruiken data uit twee bytes (16 bits) per pixel bestaat moet deze teruggebracht worden tot één byte (8-bit) per pixel.

De trainingsgegevens moeten ook aangeleverd worden in de vorm van een rasterbestand in dit formaat (*.mpd, *.mpi) met even veel rijen en kolommen als de beelden. De pixelwaarden in dit rasterbestand geven het klassennummer weer, waarbij ongetrainde pixels de waarde 0 krijgen. Zo'n rasterbestand met trainingsgegevens kan gemaakt worden in een beeldverwerkingsprogramma naar keuze. Het kan ook uit een tekstbestand vervaardigd worden met het programma *chpix* [in] [uit] [tekstbestand]. Hiervoor heb je een leeg beeld en een tekstbestand nodig. Het lege beeld kan verkregen worden met de beeldcalculator *mq* [uit]:="[beeld]*0". In het *tekstbestand* voor *chpix* moeten de beeldcoördinaten en de klassennummers van de trainingssamples staan.

1-Byte beelden kunnen afgebeeld worden met het commando *show* [beeld]. Een kleurcompositie van meerdere banden kan weergegeven worden met het commando *showcc* [beeld1] [beeld2] [beeld3]. Een rasterbestand kan afgebeeld worden met *show*, maar dan zullen de klassennummers als grijswaarden weergegeven worden. Er kan ook gebruik gemaakt worden van een legenda, door af te beelden met het commando *coldisp* [legenda] [beeld]. Hiervoor is een *legendabestand* (*.col) nodig. Dit kan gemaakt worden in een teksteditor zoals nedit. Bij zowel het gebruik van *show*, *showcc* als *coldisp* wordt het programma xv gebruikt voor het weergeven. Dit programma heeft ook mogelijkheden om de kleuren te manipuleren, in te zoomen e.d. (probeer de e-toets en de rechter muisknop).

De probabilistische segmentatie bestaat uit een integratie van classificatie en segmentatie. Omdat voor verschillende typen beelden en trainingsgegevens verschillende classificatiemethoden wenselijk zijn, is de classificatie niet in het programma voor probabilistische segmentatie opgenomen. Omdat het programma voor probabilistische segmentatie de resultaten van de classificatie nodig heeft, moet de classificatie eerst uitgevoerd worden. Bij de software zit het classificatieprogramma *annclass* [bestandenlijst-bestand].

Het programma *annclass* heeft een tekstbestandje nodig waarin de bestandsnamen gespecificeerd worden van het rasterbestand met trainingsgegevens en de beelden van de te gebruiken banden. Dit zogenaamde *bestandenlijst-bestand* (*.in) kan gemaakt worden in een teksteditor.

Annclass gebruikt het programma *ann_sample*. Voordat *ann_sample* aangeroepen kan worden, worden eerst de juiste typen invoerbestanden gemaakt. De wijze waarop dit gebeurt, is echter niet heel efficiënt. Vooral de routine *asc2annq* is erg traag bij meer dan een paar banden. Daarom zijn er voor bepaalde hoeveelheden banden (28 en 88) speciale snellere varianten gemaakt. Deze worden automatisch gebruikt wanneer een dergelijke variant aanwezig is.

Het is ook mogelijk om met een beeldverwerkingsprogramma naar keuze de classificatie uit te voeren. Voorwaarde hiervoor is dat deze methode voor iedere klasse een kans oplevert. Deze kansen zullen dan aangeleverd moeten worden als een beeld per klasse, met dezelfde naam als het *bestandenlijst-bestand* (*.in) met achtervoeging van *pd* en het klassennummer.

Na het schatten van de benodigde kansen met het classificatieprogramma kan de probabilistische segmentatie uitgevoerd worden met het programma *probseg* [bestandenlijst-bestand]. Voor *probseg* is een tweede *bestandenlijst-bestand* (*.inp) nodig met dezelfde naam als die gebruikt is voor het classificatieprogramma. In dit tekstbestand worden de namen gespecificeerd van het rasterbestand voor validatie en van de banden die gebruikt worden voor de segmentatie. Enkele instellingen voor het programma *probseg*, zoals de drempelwaarden voor de segmentatiepiramide, kunnen gewijzigd worden in het script *probseg* met behulp van een teksteditor. Van belang is dat deze drempelwaarden zoveel mogelijk in het interval gekozen worden, waarbinnen voor het betreffende beeld niet te veel of te weinig segmenten gevormd worden.

Het programma *preseg* [naam] [segmentatie] [test] doet hetzelfde als het programma *probseg*. Het verschil is dat dit *preseg* een zelf op te geven segmentatie gebruikt. Een segmentatiebestand is een rasterbestand (*.mpd, *.mpi) waarin alle pixelwaarden een segmentnummer weergeven. Ieder segment moet een uniek nummer hebben. Om meer dan 256 segmenten aan te kunnen, hebben de segmentbestanden 4 bytes per pixel. Zo'n segmentatiebestand kan gemaakt worden uit een 1-byte-bestand zonder unieke nummers met het commando *makeseg* [beeld] [segmentatie].

Enkele opmerkingen over de programma's *annclass* en *probseg*:

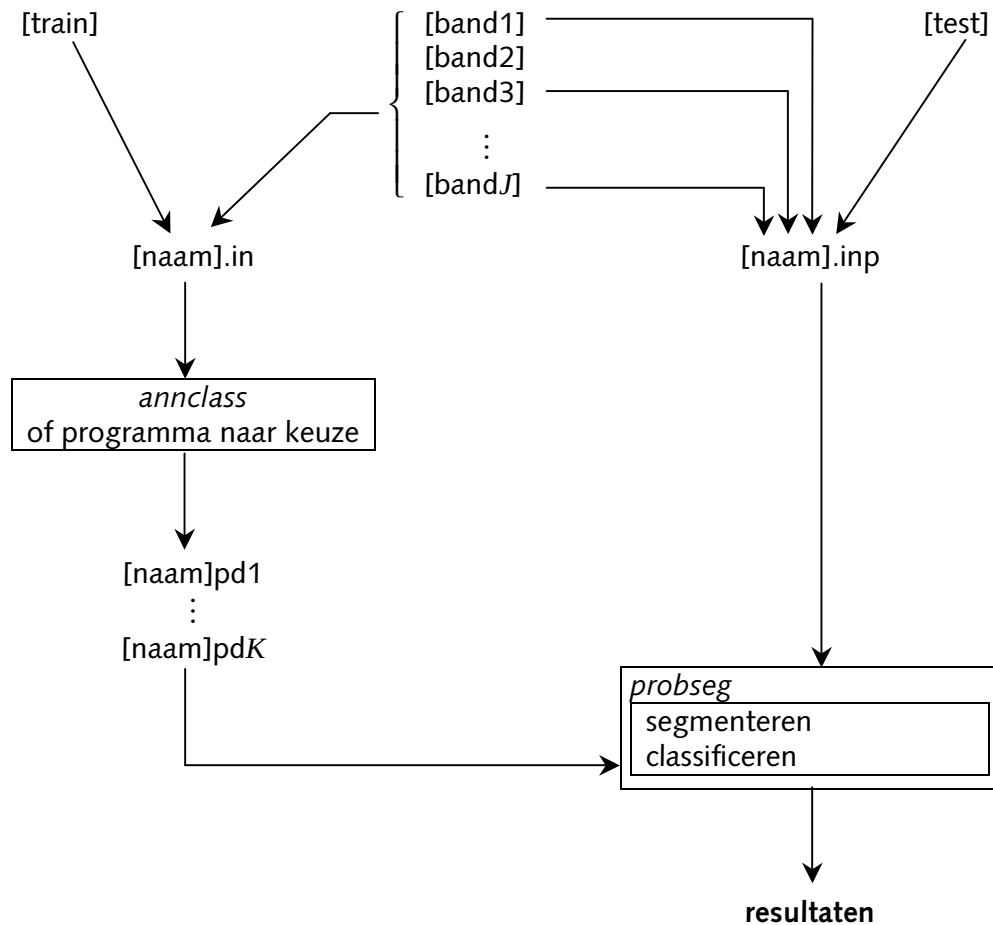
- Alle normale uitvoer op het scherm begint met een verticale streep (|). Iedere andere mededeling op het scherm duidt er op dat er iets mis gaat, of dat bij één van de eerdere delen van het programma iets mis gegaan is.
- Het programma kan beter niet tegelijkertijd meerdere keren uitgevoerd worden. Voor de meeste delen van het programma is dit geen probleem, maar bepaalde delen zouden (tijdelijke) bestanden met dezelfde naam kunnen gaan maken waarbij de één de ander overschrijft.
- Het programma heeft veel vrije schijfruimte nodig voor tijdelijke bestanden die aan het eind van het programma weer verwijderd worden. Dit kan oplopen tot 20 Gigabyte voor 5,5 miljoen pixels, 88 banden en 21 klassen.
- De rekentijd kan aanzienlijk zijn, van ca. 4 minuten voor een beeld met 75 duizend pixels, 3 banden en 7 klassen tot ruim 14 uur voor 5,5 miljoen pixels, 88 banden en 21 klassen.
- Er zijn verschillende beperkingen voor het gebruik van het programma voor grote hoeveelheid data. Eén reden is dat er geen (tijdelijke) bestanden groter dan 2 Gigabyte gemaakt mogen worden. Op een 32-bit-systeem kunnen bestanden groter dan 2^{31} bytes namelijk niet door alle programma's verwerkt worden. *Annclass* kan daarom alleen data aan waarvoor geldt:

$$N \cdot (4 \cdot J + 1) \leq 2^{31}$$

met: N het aantal pixels
 J het aantal banden

Een andere reden is dat in sommige routines een maximum aantal segmenten, klassen of banden gedeclareerd is. Ook om andere redenen kunnen er beperkingen zijn voor bepaalde routines. De maxima waar het programma tot op heden voor gebruikt is zijn:

5 531 648 pixels
 2 584 259 segmenten
 21 klassen
 127 banden



Figuur B2.1. Schema van het gebruik van de programmatuur. Met [naam] de naam van de bestandenlijst-bestanden.

Resultaten

De resultaten van de probabilistische segmentatie worden opgeslagen in de map waar het commando gegeven is. Dit zijn zes classificaties: een crisp en fuzzy classificatie op grond van kansdichtheden (pd), a priori kansen per segment (prior) en a posteriori kansen (post). De crisp classificaties bestaan uit een rasterbestand met het klassennummer als pixelwaarde. De fuzzy classificaties bestaan uit een rasterbestand per klasse. In dit rasterbestand per klasse geven de pixelwaarden de kans op deze klasse weer (vermenigvuldigd met 255). Daarnaast worden nog tekstbestanden met confusionmatrices gemaakt evenals enkele segmentatiebestanden.

Onderstaande lijst beschrijft de inhoud van de bestanden die het resultaat zijn van de probabilistische segmentatie. Hierin is [naam] de bestandsnaam van de *bestandenlijst-bestanden* (*.in en *.inp).

[naam]pd	Crisp classificatie op grond van de maximale kans per pixel zonder probabilistische segmentatie
[naam]prior	Crisp classificatie op grond van de maximale a priori kans per segment
[naam]post.	Crisp classificatie op grond van de maximale a posteriori kans per pixel

[naam]pd[nr.]	Fuzzy classificatie: rasterbestand met de conditionele kans op een featurevector per pixel zonder probabilistische segmentatie voor de klasse [nr.]
[naam]prior[nr.]	Fuzzy classificatie: rasterbestand met a priori kans per segment voor de klasse [nr.]
[naam]post[nr.]	Fuzzy classificatie: rasterbestand met a posteriori kans per pixel voor de klasse [nr.]
[naam]seglvl[nr.]	Segmentatiebestand van piramideniveau [nr.]
[naam]seg	Segmentatiebestand van de geselecteerde segmenten
[naam]numberofclasses	Rasterbestand met het aantal klassen per segment (na verwaarlozing van kleine klassenaandelen)
[naam]pd.mat	Tekstbestand: confusionmatrix van [naam]pd
[naam]prior.mat	Tekstbestand: confusionmatrix van [naam]prior
[naam]post.mat	Tekstbestand: confusionmatrix van [naam]post

Met de beeldcalculator *mq* (alleen geschikt voor 1-byte bestanden) kan desgewenst extra uitvoer gemaakt worden, zoals een crisp classificatie van alleen de segmenten met maar één klasse. Een confusionmatrix van een classificatie kan bepaald worden met het programma *compare* [test] [class]. De segmentatiebestanden kunnen weergegeven worden met het commando *showseg* [segmentatie]. Hiervoor wordt een 1-byteversie van het segmentatiebestand gemaakt, waarin segmentnummers niet meer uniek zijn. N.B.: het kan gebeuren dat twee aangrenzende segmenten dan hetzelfde nummer krijgen.

Commando's

Hieronder staan beschrijvingen van de verschillende commando's. Tenzij expliciet vermeld moeten bestandsnamen zonder extensie opgegeven worden.

annclass [bestandenlijst-bestand] [k]
 [bestandenlijst-bestand] de naam van het bestandenlijst-bestand voor classificatie (*.in)
 [k] het aantal nearest neighbours

chpix [in] [uit] [tekstbestand]
 [in] de naam van het te wijzigen rasterbestand (*.mpd, *.mpi)
 [uit] de naam van het te maken rasterbestand (*.mpd, *.mpi)
 [tekstbestand] de naam van het tekstbestand voor chpix inclusief extensie

coldisp [legenda] [beeld]
 [legenda] de naam van het legendabestand (*.col)
 [beeld] de bestandsnaam van het weer te geven beeld (*.mpd, *.mpi)

compare [train] [class]
 [train] naam van het rasterbestand met trainingsgegevens
 [class] naam van het rasterbestand met de classificatie

makempi [naam] [#kol] [#rij] [#byte]
 [naam] de naam van het te maken mpi-bestand
 [#kol] het aantal kolommen
 [#rij] het aantal rijen
 [#byte] het aantal bytes per pixel

makeseg [beeld] [segmentatie]

[beeld] naam van bestand voor segmentatie (*.mpd, *.mpi)

[segmentatie] naam van het te maken segmentatiebestand

mpi [naam] nl nc.

[naam] de naam van het mpi-bestand

mq [uit]:="[bewerking ([beeld1], [beeld2])]"

[uit] naam van het te maken bestand (*.mpd, *.mpi)

[beeld1] en [beeld2] de bestandsnamen van de te gebruiken beelden (1-byte *.mpd, *.mpi)

Voorbeelden:

mq [uit]:="[beeld1]*[beeld2]/255"

mq [uit]:="[if([beeld1]=[beeld2],255,0)"]

preseg [naam] [segmentatie] [test]

[naam] naam van de bestanden gemaakt door annclass (exclusief pd)

[segmentatie] naam van het segmentatiebestand (*.mpd, *.mpi)

[test] de naam van het rasterbestand voor validatie (*.mpd, *.mpi)

probseg [bestandenlijst-bestand]

[bestandenlijst-bestand] de naam van het bestandenlijst-bestand voor probabilistische segmentatie (*.inp)

show [beeld]

[beeld] de bestandsnaam van het weer te geven beeld (*.mpd, *.mpi)

showcc [band1] [band2] [band3]

[band1], [band2] en [band3] de bestandsnamen van de weer te geven banden (*.mpd, *.mpi)

showseg [segmentatie]

[segmentatie] segmentatiebestand

Tekstbestanden

Legendabestand (*.col)

red%	green%	blue%	#0
0	0	0	
[R1]	[G1]	[B1]	
⋮	⋮	⋮	
[RK]	[GK]	[BK]	

[R1] ... [RK] de roodwaarden tussen 0 en 1000 voor klasse 1 tot en met K

[G1] ... [GK] de groenwaarden tussen 0 en 1000 voor klasse 1 tot en met K

[B1] ... [BK] de blauwwaarden tussen 0 en 1000 voor klasse 1 tot en met K

Tekstbestand voor chpix (willekeurige extensie)

```
[c1] [r1] [v1]  
⋮   ⋮   ⋮  
[cN] [rN] [vN]
```

[c1] ... [cN] de kolomnummers van de te wijzigen pixels 1 tot en met N
[r1] ... [rN] de rijnummers van de te wijzigen pixels 1 tot en met N
[v1] ... [vN] de pixelwaarden van de te wijzigen pixels 1 tot en met N
N.B.: Het pixel in de linker bovenhoek heeft kolom- en rijnummer 0

Bestandenlijst-bestand voor classificatie (*.in)

```
[train]  
[band1]  
⋮  
[bandJ]
```

[train] de naam van het rasterbestand met trainingsgegevens (*.mpd, *.mpi)
[band1] ... [bandJ] de bestandsnamen van de banden (*.mpd, *.mpi)

Bestandenlijst-bestand voor probabilistische segmentatie (*.inp)

```
[test]  
[band1]  
[band2]  
[band3]
```

[test] de naam van het rasterbestand voor validatie (*.mpd, *.mpi)
[band1], [band2] en [band3] de bestandsnamen van de drie banden (*.mpd, *.mpi)

Symbolen

Onderstaande lijst bevat alle symbolen die niet bij ieder gebruik toegelicht worden.

$\dots^o$	Benaderde waarde voor ...
A	Matrix, verband tussen Δx en Δy
$A()$	Functie, verband tussen x en y
b_j	Band j
C_k	Klasse k
d	Drempelwaarde waaruit twee andere drempelwaarden d_1 en d_2 voor de segmentatie berekend worden
$f_{C_k s_i}$	Fractie van klasse k in sample i
J	Aantal banden
K	Aantal klassen
k	Aantal neighbours voor k-nearest-neighbours-classificatie
$P_A^\perp$	$I - A \cdot (A^T Q_y^{-1} A)^{-1} A^T Q_y^{-1}$
$P(C_k)$	Kans dat een pixel tot klasse k behoort, ongeacht de featurevector van het pixel (a priori kans op een klasse)
$P(C_k \chi_p)$	Kans dat een pixel p met featurevector χ_p tot klasse k behoort (a posteriori kans op een klasse)
$P(\chi_p)$	Kans dat een pixel een featurevector χ_p heeft
$P(\chi_p C_k)$	Kans dat een pixel p van klasse k een featurevector χ_p heeft (conditionele kans op een featurevector)
Q	(Co)variantiematrix
s_i	Sample i
x	Vector met onbekenden
$\hat{x}$	Kleinstekwadratenoplossing voor x
y	Vector met waarnemingen
$\hat{\sigma}$	Kleinstekwadratenoplossing voor de vector met (co)variantiefactoren
$\chi_{b_j C_k}$	Gemiddelde grijswaarde voor band j van klasse k
$\chi_{b_j s_i}$	Grijswaarde voor band j van sample i
χ_{C_k}	Gemiddelde featurevector van klasse k
χ_p	Featurevector van pixel p



HyMap-beeld Schiermonnikoog stroken 2 en 4