

Verifikation von Ackerland- und Grünlandobjekten eines topographischen Datensatzes mit monotemporalen Bildern

Von der Fakultät für Bauingenieurwesen und Geodäsie
der Gottfried Wilhelm Leibniz Universität Hannover
zur Erlangung des Grades

DOKTOR-INGENIEUR (Dr.-Ing.)

genehmigte Dissertation
von

Dipl.-Ing. Petra Helmholz

geboren am 25.04.1979 in Halberstadt

Vorsitzender der Prüfungskommission:	Prof. Dr.-Ing. habil Jürgen Müller
Referenten:	Prof. Dr.-Ing. habil Christian Heipke Prof. Dr.-Ing. Markus Gerke PD Dr.techn. Franz Rotternsteiner
Gutachter:	Prof. Dr. Clive Fraser

Tag der mündlichen Prüfung: 19. Januar 2012.

Die Dissertation wurde am 25.10.2012 bei der
Fakultät für Bauingenieurwesen und Geodäsie der Gottfried Wilhelm Leibniz Universität Hannover
eingereicht.

Kurzfassung

In dieser Arbeit wird ein neuer Ansatz zur automatischen Qualitätsbewertung bestehender 2D-Vektordaten vorgestellt, die in Geoinformationssystemen (GIS) gespeichert sind. Die notwendigen Informationen für die Qualitätskontrolle werden mit Hilfe von Bildanalyseverfahren aus aktuellen Fernerkundungsdaten gewonnen. Im Gegensatz zu anderen Verfahren liegt der Schwerpunkt in dieser Arbeit auf der Qualitätsbewertung von GIS-*Ackerland*- und *Grünland*objekten, die von besonderem Interesse sind, da sie wesentlich zur Nahrungsversorgung der Bevölkerung beitragen. *Acker*- und *Grünland* bedecken gleichzeitig den größten Anteil an der Gesamtfläche Deutschlands mit über 50%. Diese Arbeit fokussiert auf die Verifikation als Teil der Qualitätsbewertung.

Das vorgestellte semi-automatische System hat die Verifikation unter Verwendung orthorektifizierter monotemporaler Luft- und Satellitenbilder mit einer geometrischen Auflösung von 0,5 bis 1 m zum Ziel. Spezialisiert ist das System auf GIS, die dem Inhalt und Detaillierungsgrad einer topographischen Karte mittleren Maßstabs entsprechen. Ziel der Verifikation ist es, den Anteil der inkorrekten Objekte im GIS auf höchstens 5% zu senken. Zeitgleich soll der manuelle Aufwand, der für den menschlichen Bearbeiter zur Qualitätskontrolle notwendig ist, durch die Verwendung des Systems minimiert werden. Dies wird erreicht, indem ein Teil der GIS-Objekte in einem Verifikationsverfahren automatisch verifiziert wird und daher vom Bearbeiter nicht mehr betrachtet werden muss.

Das entwickelte Verifikationsverfahren ist zweistufig. Zunächst wird im ersten Schritt die gemeinsame Klasse *Acker-/Grünland* von anderen Klassen wie *Siedlung*, *Industrie* und *Wald* mittels einer pixelbasierten Klassifikation getrennt. Wird ein GIS-*Acker-/Grünland*-Objekt einer anderen Klasse zugewiesen, so wird dieses vom System abgelehnt und muss durch den menschlichen Bearbeiter manuell verifiziert werden. Im zweiten Schritt werden alle anderen GIS-*Ackerland*- und *Grünland*objekte mittels einer objektbasierten Klassifikation entweder der Klasse *Ackerland* oder *Grünland* zugewiesen. Ist eine Klassifikation nicht möglich, wird das GIS-Objekt einer *Zurückweisungsklasse* zugewiesen. Durch die Kombination eines pixel- und eines objektbasierten Ansatzes sollen die Vorteile beider Verfahren genutzt werden. Nachdem alle GIS-Objekte einer Klasse zugewiesen wurden, wird das Ergebnis mit der Eintragung im GIS verglichen. Stimmt die ermittelte Klasse mit der im GIS überein, wird das GIS-Objekt als richtig angenommen, wenn nicht, wird es verworfen, und eine manuelle Kontrolle des GIS-Objektes durch den menschlichen Bearbeiter ist notwendig.

Der Schwerpunkt der Arbeit liegt auf dem zweiten Schritt. Da in GIS-*Ackerland*- und *Grünland*objekten häufig mehrere Bewirtschaftungseinheiten (Schläge) vorhanden sind, wird zunächst eine Segmentierung mittels Wasserscheidentransformation durchgeführt. Ziel ist, dass ein Segment einem Schlag entspricht. Die Zugehörigkeit eines Segmentes zur Klasse *Ackerland* oder *Grünland* wird anschließend durch eine objektbasierte Klassifikation mittels spektraler, textueller, struktureller und geometrischer Merkmale bestimmt. Als Klassifikator dient das Verfahren der *Support Vector Machines*. Alle Segmente, die zu klein sind, um klassifiziert zu werden, werden einer *Zurückweisungsklasse* zugewiesen. Sind alle Segmente eines GIS-Objektes einer Klasse zugewiesen, wird die Flächensumme je Klasse berechnet. Ist die Flächensumme der *Zurückweisungsklasse* größer als die Summe der Flächen, die klassifiziert werden konnten, gilt das GIS-Objekt als nicht klassifizierbar und wird der *Zurückweisungsklasse* zugeordnet. Anderenfalls wird das GIS-Objekt der Klasse zugeordnet, deren Flächensumme die Mindestkartierfläche überschreitet. Wird von beiden oder von keiner Klasse die Mindestkartierfläche erreicht, wird das GIS-Objekt ebenfalls der *Zurückweisungsklasse* zugeordnet.

Das automatische Verifikationsverfahren wurde unter verschiedenen Gesichtspunkten an Hand von Beispielen des ATKIS Basis-DLM und eines Schlagkatasters getestet, die insgesamt 3313 GIS-*Ackerland*- und *Grünland*objekte beinhalten und eine Fläche von über 302 km² bedecken. Zunächst wurde bei der separaten Betrachtung des ersten Schrittes festgestellt, dass dieser für die Verifikation von GIS-*Ackerland*- und *Grünland*objekten geeignet ist. Ein ähnliches Ergebnis wurde bei der separaten Betrachtung des zweiten Schrittes erreicht. Nähere Untersuchungen des zweiten Schrittes ergaben, dass der Erfolg der Verifikation vom Zeitpunkt der Aufnahme des Bildes abhängt. Bei der Betrachtung der Merkmalsgruppen für die objektbasierte Klassifikation wurde festgestellt, dass die besten Ergebnisse immer dann erreicht werden können, wenn spektrale und strukturelle Merkmale beteiligt sind. Insgesamt konnte bei der Kombination beider Schritte eine erfolgreiche Verifikation von GIS-*Ackerland*- und *Grünland*objekten durchgeführt werden. Der Anteil der inkorrekten GIS-Objekte nach dem Verifikationsprozess beträgt weniger als 5% und der Bearbeiter spart mindestens die Hälfte der Zeit bei der Qualitätskontrolle unter Verwendung des Systems verglichen zum Arbeitsaufwand ohne Verwendung des Verifikationssystems.

Stichworte: Verifikation, 2D-Vektordaten, monotemporal, Bildanalyse, Ackerland, Grünland

Abstract

In this thesis a new approach for the automatic quality assessment of 2D vector data is introduced. 2D vector data are managed in geo-information systems (GIS). The information needed for the quality assessment is extracted from up-to-date remote sensing data using image analysis techniques. In contrast to other approaches the main goal is the quality assessment of *cropland* and *grassland* GIS-objects, which are of interest as they assure the food supply of the population, and at the same time cover more than 50% of the area of Germany. This thesis focuses on the verification as a part of the quality assessment.

The goal of the introduced semi-automated method is the verification of *cropland* and *grassland* GIS-objects using orthorectified mono-temporal aerial and satellite images with a geometric resolution between 0.5 to 1 m. The method is specialised in GIS that contain information of a medium scale topographic map. The goal of the verification process is to reduce the percentage of incorrect objects in the GIS to a maximum of 5%. At the same time, the manual effort that is required by a human operator for the quality assessment has to be reduced to a minimum using the system. This is achieved because GIS-objects which have been accepted by the automatic verification method do not have to be reviewed by the human operator.

The developed verification approach consists of two analysing steps. In the first analysing step, the joined class *cropland/grassland* (also called *agriculture*) is separated from other classes such as *settlement*, *industry* and *forest* using a pixel-based classification algorithm. If a *cropland* or *grassland* GIS-object was not assigned to *agriculture*, this GIS-object will be rejected by the system and has to be reviewed by a human operator. In the second analysing step, all remaining *cropland* and *grassland* GIS-objects are separated into the two classes *cropland* and *grassland* using an object-based classification algorithm. If a classification is not possible, the GIS-object is assigned to a *rejection class*. Due to the combination of a pixel-based and an object-based classification algorithm, the advantages of both strategies are exploited. After all GIS-objects were assigned to a class, the classification result is compared to the original object classes in the GIS. If the detected class and the original class in the GIS are equal, the GIS-object is accepted by the system, and a manual check of this GIS-object is not necessary.

This thesis focuses on the second analysing step. Due to the fact that *cropland* and *grassland* GIS-objects often contain different management units, segmentation is necessary. The goal of the segmentation is that one management unit corresponds to one segment. Whether a segment belongs to *cropland* or *grassland* is determined using spectral, textural, structural and geometric features in an object-based classification. The algorithm of the *Support Vector Machines* is used for this step. All segments that are too small for the object-based classification are assigned to the *rejection class*. After all segments of a GIS-object have been classified, the area sum per class is calculated. If the sum of the areas of the *rejection class* is larger than the sum of the areas of all other segments, this GIS-object cannot be classified and hence is assigned to the *rejection class*. Otherwise, the GIS-object is assigned to the class whose area sum is larger than the minimum mapping area. If both or none of the classes achieve the minimum mapping area, the GIS-object is also assigned to the *rejection class*.

The automatic verification approach was evaluated considering ATKIS Basis-DLM (German Authoritative Topographic Cartographic Information System) and a local thematic dataset where a GIS-object contains only one management unit, with in total 3313 GIS-*cropland*- and *grassland* objects, covering a total area of 302 km². The evaluation of the first separately tested analysis step could prove that it is suitable for the verification task. A similar result was achieved by evaluating the second analysis step separately. In addition, further analyses of the second analysis step have been done with two main results. First, the success of the verification depends on the time when the image was acquired. Second, the best classification and verification results were achieved using spectral and structural features. By combining both analysis steps, a successful verification of GIS-*cropland*- and *grassland* objects was possible. The percentage of incorrect GIS-objects could be reduced to less than 5%, and at the same time the human operator saves at least half of the time for the quality assessment by using the proposed method compared to a complete manual procedure.

Keywords: verification, 2D vector data, mono-temporal, image analysis, *cropland*, *grassland*

Inhaltsverzeichnis

1. Einleitung.....	1
1.1. Qualität von Geodaten.....	1
1.2. Verifikation landwirtschaftlicher Nutzflächen.....	2
1.3. Ziel und wissenschaftlicher Beitrag der Arbeit.....	2
1.4. Struktur der Arbeit	3
2. Stand der Wissenschaft.....	5
2.1. Klassifikation von landwirtschaftlichen Nutzflächen	5
2.1.1. Klassifikationsalgorithmen	6
2.1.2. Merkmale für die Klassifikation	9
2.1.3. Beispiele für die Klassifikation landwirtschaftlicher Flächen	12
2.2. Verifikation von GIS-Objekten	18
2.3. Segmentierung von Schlägen in GIS-Objekten	20
2.4. Diskussion	21
3. Grundlagen.....	23
3.1. Einführung in den Klassifikationsalgorithmus der Support Vector Machine	23
3.1.1. Funktionsweise einer SVM.....	23
3.1.2. SVM mit Fehlern in den Trainingsdaten	27
3.1.3. Anmerkungen zur praktischen Durchführung	28
3.1.4. Automatische Bestimmung der Parameter der SVM.....	29
3.1.5. Anwenden einer SVM auf ein Mehrklassenproblem.....	30
3.1.6. Vergleich mit anderen Klassifikationsalgorithmen	30
3.1.7. Zusammenfassung.....	31
3.2. Einführung in die Segmentierung mit Wasserscheidentransformation und Regionennachbarschaftsgraph.....	31
4. Neue Methode zur Verifikation der Objektarten Ackerland und Grünland eines GIS.....	33
4.1. Eingangsdaten	33
4.2. Überblick.....	33
4.3. Klassifikation zur Identifizierung von GIS-Ackerland- und Grünlandobjekten	35
4.3.1. Pixelbasierte Klassifikation	36
4.3.2. Überführung der pixelbasierten Klassifikation auf ein GIS-Objekt	37
4.4. Segmentierung von Schlägen	38
4.4.1. Attribute der Bildregionen und deren Grenzen.....	39
4.4.2. Verschmelzen von Regionen	39
4.4.3. Nachbearbeitung des Segmentierungsergebnisses an Hand von GIS-Spezifikationen	42
4.5. Klassifikation eines Segmentes	42
4.5.1. Berechnung der Merkmale für die Klassifikation.....	42
4.5.2. Klassifikation anhand der Merkmale	51
4.6. Klassifikation und Verifikation des GIS Objektes	52
4.7. Nachbearbeitung.....	52
4.8. Grenzen des Verfahrens	52
4.8.1. Pixelbasierte Klassifikation zur Identifizierung von Ackerland- / Grünlandobjekten.....	53
4.8.2. Segmentierung von Schlägen.....	53
4.8.3. Klassifikation eines Segmentes.....	54
4.9. Zusammenfassung	56

5.	<i>Evaluation</i>	57
5.1.	Bewertung der Kriterien für die Klassifikations- und Verifikationsergebnisse.....	57
5.1.1.	Bewertung von Klassifikations- und Verifikationsergebnissen.....	57
5.1.2.	Kriterien für die Verifikationsergebnisse.....	59
5.2.	Datengrundlage	59
5.2.1.	GIS-Datensätze	59
5.2.2.	Bilddaten	60
5.2.3.	Testszenen.....	61
5.3.	Wahl der Parameter	66
5.3.1.	Parameter für die Klassifikation zur Identifizierung von GIS-Ackerland- und Grünlandobjekten.....	67
5.3.2.	Parameter für die Segmentierung.....	68
5.3.3.	Parameter für die Klassifikation eines Segmentes.....	68
5.3.4.	Parameter für die Klassifikation und Verifikation eines GIS-Objektes.....	69
5.4.	Evaluation des ersten Analyseschrittes	69
5.4.1.	Untersuchung der Abhängigkeit des Verifikationsergebnisses von s_q	70
5.4.2.	Untersuchung des Erfolgs der Verifikation für zwei verschiedene s_q in unterschiedlichen Testszenen.....	71
5.4.3.	Bewertung	73
5.5.	Untersuchung des Einflusses von verschiedenen Merkmalsgruppen auf die objektbasierte Klassifikation	74
5.5.1.	Evaluierung der Klassifikationsergebnisse	75
5.5.2.	Evaluierung der Verifikationsergebnisse	76
5.5.3.	Bewertung	78
5.6.	Untersuchung des Einflusses von verschiedenen Zeitpunkten auf die objektbasierte Klassifikation	79
5.6.1.	Evaluierung der Klassifikationsergebnisse	79
5.6.2.	Evaluierung der Verifikationsergebnisse	80
5.6.3.	Bewertung	83
5.7.	Evaluation des zweiten Analyseschrittes	83
5.7.1.	Evaluierung der Klassifikationsergebnisse	84
5.7.2.	Evaluierung der Verifikationsergebnisse	88
5.7.3.	Bewertung	93
5.8.	Evaluierung des gesamten Algorithmus.....	94
5.8.1.	Evaluierung der Verifikationsergebnisse	94
5.8.2.	Bewertung	99
5.9.	Zusammenfassung.....	99
6.	<i>Folgerungen und Ausblick</i>	101
	<i>Literaturverzeichnis</i>	103
	<i>Abkürzungsverzeichnis</i>	110
	<i>Lebenslauf/ Curriculum vitae</i>	111
	<i>Dank</i>	112

1. Einleitung

1.1. Qualität von Geodaten

Viele Entscheidungen des öffentlichen und privaten Lebens basieren auf raumbezogenen Daten, die in Geoinformationssystemen (GIS) gespeichert und gepflegt werden. Die Brauchbarkeit der Daten für die Anwendung hängt von deren Qualität ab. Unter Qualität versteht man den „Grad, in dem ein Satz inhärenter Merkmale Anforderungen erfüllt“ (ISO 9000, 2005, S. 18). Aufgabe eines GIS ist die abstrakte Repräsentation der realen Welt in einer Datenbank. Um die reale Welt in einem GIS repräsentieren zu können, sind Abbildungsvorschriften, die das abstrakte Abbild der realen Welt definieren, notwendig. Diese Abbildungsvorschriften werden in GIS-Objektartenkatalogen zusammengefasst. Der kleinste Bestandteil eines GIS ist ein GIS-Objekt, das einer GIS-Objektart (Objektklasse) zugeordnet ist, welche im GIS-Objektartenkatalog beschrieben ist. Soll ein GIS die oben genannte Qualitätsdefinition nach ISO erfüllen, so muss der GIS-Objektartenkatalog die reale Welt widerspruchsfrei repräsentieren (Modellqualität) und die Daten müssen ihren Spezifikationen entsprechen (Datenqualität) (Joos, 2000). Es gibt 5 wichtige Maße für die Qualität von Geodaten: *Vollständigkeit* (*completeness*), *logische Konsistenz* (*logical consistency*), *geometrische Genauigkeit* (*positional accuracy*), *zeitliche Genauigkeit* (*temporal accuracy*) und *thematische Genauigkeit* (*thematic accuracy*) (ISO 19113, 1999 und ENV 12656, CEN 1998).

Nur die Konsistenz der Daten kann automatisch ohne einen Vergleich mit der realen Welt kontrolliert werden, indem ein gegebenes Datenmodell mit den vorliegenden GIS-Daten verglichen wird. Viele Aspekte der anderen Qualitätsmaße (*Vollständigkeit*, *geometrische*-, *zeitliche*- und *thematische Genauigkeit*) können nur durch Vergleich der GIS-Daten mit der realen Welt bewertet werden. Die reale Welt kann z.B. in Form von aktuelleren Luft- und Satellitenbildern repräsentiert sein. Die aktuellen Luft- und Satellitenbilder bilden dann die Referenz.

Der Vergleich der GIS-Daten mit der Referenz zum Finden aller Abweichungen, d.h. das Finden von Objekten im GIS, die in der Realität nicht existieren, oder die in den GIS-Daten falsch repräsentiert sind, wird *Verifikation* genannt. Die Verifikation ist Teil des Qualitätsmanagements von GIS-Daten, wobei nicht alle GIS-Objektarten mittels Luft- und Satellitenbildern überprüft werden können. Bei einer Verifikation von GIS-Objekten bezüglich der *thematischen Genauigkeit* unter Verwendung aktueller Bilddaten wird zeitgleich auch die *zeitliche Genauigkeit* überprüft. Der Überprüfung der *thematischen Genauigkeit* von GIS-Objekten kommt damit eine besondere Bedeutung zu. Bei der Verifikation ist die Ursache der Abweichung ohne Bedeutung. Ein Verfahren, das speziell *zeitliche Genauigkeit* untersucht, ist die *Veränderungsdetektion* (*change detection*). Ziel der *change detection* ist das Erkennen von Änderungen zwischen zwei Zeitpunkten t_1 und t_2 , wobei zwei Varianten existieren. Zum einen kann eine *change detection* auf Grundlage zweier Bilder, die zu den Zeitpunkten t_1 und t_2 aufgenommen wurden, oder zum anderen auf Grundlage eines GIS von Zeitpunkt t_1 mit einem aktuelleren Bild zum Zeitpunkt t_2 erfolgen. Auch eine Kombination beider Ansätze ist möglich. Die Verifikation von flächenhaften GIS-Datensätzen ohne Überlagerungsflächen¹ und ohne Lücken schließt eine *change detection* in dem Sinn ein, dass es eine Veränderung gegeben hat, aber nicht, dass man die Art der Veränderung erkennt. Die Aktualisierung (*update*) hingegen umfasst sowohl die Ergebnisse einer Verifikation also auch die Erfassung neuer Objekte, wobei die Übernahme der Änderungen in ein GIS ebenfalls dazugehört. Eine detailliertere Diskussion der Terminologie ist (Gerke und Heipke, 2008) zu entnehmen.

Gegenwärtig erfolgt die Qualitätskontrolle von Geodaten in der Praxis meist manuell (Krickel, 2010). Um den Aufwand der manuellen Arbeit zu reduzieren, ist ein hoher Grad der Automatisierung notwendig. Die automatisierte Qualitätskontrolle von GIS-Datenbanken rückte daher mit der Einführung neuer hochauflösender digitaler bildgebender Sensoren in den Fokus der Forschung (Hoffmann et al., 2000), (Alpin et al., 1999b).

Die vorliegende Arbeit beschäftigt sich mit der Automatisierung der Verifikation der *thematischen Genauigkeit* der beiden GIS-Objektarten Acker- und Grünland im Rahmen der Qualitätskontrolle von GIS-Datenbanken.

¹ Überlagerungsflächen sind Flächen einer Objektart, die von mindestens einer Fläche einer anderen Objektart überlagert wird.

1.2. Verifikation landwirtschaftlicher Nutzflächen

Bei der Verifikation von GIS-Datenbanken sind Objektarten von besonderem Interesse, die volkswirtschaftlich relevant sind. Dazu gehören insbesondere landwirtschaftliche Flächen wie Acker- und Grünland, da diese für die Nahrungsversorgung der Bevölkerung wichtig sind. In Deutschland bedecken landwirtschaftliche Nutzflächen über 50% der Gesamtfläche des Territoriums. Der Anteil von Ackerland nimmt dabei ca. 70% und von Grünland, der zweitgrößte Anteil, ca. 29% der gesamten landwirtschaftlichen Nutzflächen ein (Statistisches Bundesamt, 2011). Die automatisierte Verifikation von GIS-Objekten der Objektarten Acker- und Grünland ist daher von besonderem Interesse.

Die automatisierte Verifikation von GIS-Ackerland- und Grünlandobjekten ist jedoch sehr komplex. Zum einen können die GIS-Objekte der Objektklasse Ackerland sehr unterschiedliche Eigenschaften annehmen. So kann ein GIS-Ackerlandobjekt mit Vegetation bedeckt (bewachsen) oder nicht bedeckt (nicht bewachsen) sein oder auch von unterschiedlichen Fruchtarten bedeckt sein. Zum anderen gibt es viele Spezifikationen, die durch den GIS-Objektenkatalog vorgegeben sind. Diese GIS-Spezifikationen sind vom Detaillierungsgrad des jeweiligen GIS abhängig. Der Inhalt und damit der Detaillierungsgrad des GIS orientiert sich dabei am Inhalt einer topographischen Karte mit den entsprechenden Kartenmaßstäben. Eine übliche Spezifikation für viele GIS, die den Inhalt und Detaillierungsgrad einer topographischen Karte mittleren Maßstabs entsprechen, ist z.B., dass Ackerland- oder Grünlandobjekte mehr als eine Bewirtschaftungseinheit beinhalten. Eine Bewirtschaftungseinheit wird auch als Schlag bezeichnet. Besteht ein GIS-Objekt aus nur einer Bewirtschaftungseinheit, also nur einem Schlag, so wird dieses GIS auch als Schlagkataster bezeichnet. Die meisten GIS sind jedoch keine Schlagkataster. Dies ist insbesondere beim deutschen *Amtlichen Topographisch-Kartographischen Informationssystem - ATKIS* (AdV, 1997), dem europäischen *CORINE Land Cover - CLC* (EEA, 2011) oder dem internationalen *Multinational Geospatial Co-production Program - MGCP* der Fall. Eine weitere übliche Spezifikation von verschiedenen GIS ist, dass in Ackerland- und Grünlandobjekten bestimmte andere Objektarten zugelassen sind (z.B. Wald oder eine Grünlandfläche in einem GIS-Ackerlandobjekt bzw. umgekehrt), solange die zugehörige Fläche unter der im Objektenkatalog festgeschriebenen Mindestkartierfläche liegt. Andere Objektarten hingegen sind ausgeschlossen. Dazu gehören z.B. in ATKIS Siedlungsfläche oder einzelne Häuser als Teil einer Siedlungsfläche.

Möchte man GIS-Objekte der Objektklassen Acker- und Grünland verifizieren, muss mit diesen Spezifikationen sowie mit den unterschiedlichen Erscheinungsformen von Ackerland- und Grünlandobjekten umgegangen werden. Die Verifikation stellt damit eine komplexere Aufgabe als eine Klassifikation dar. Bei der Klassifikation werden Objekte nur an Hand ihrer Merkmale einer Klasse zugewiesen. Spezifikationen spielen bei der Zuordnung eines Objektes zu einer Klasse explizit keine Rolle.

1.3. Ziel und wissenschaftlicher Beitrag der Arbeit

Ziel der vorliegenden Arbeit ist die Entwicklung einer Methode für die automatische Verifikation von GIS-Objekten der Objektarten Acker- und Grünland unter Verwendung orthorektifizierter monotemporaler Luft- und Satellitenbilder mit einer geometrischen Auflösung von 0,5 bis 1 m. Bei einer komplett automatischen Kontrolle ist die Zeitersparnis, die ein menschlicher Bearbeiter unter Verwendung des Systems erreichen kann, maximiert. Erfahrungen der Autorin zeigen, dass ein GIS im Normalfall nur wenige Fehler enthält und daher die Wahrscheinlichkeit, dass genügend viele Fehler gefunden werden, um eine gewünschte *thematische Genauigkeit* des GIS zu erreichen, bei einer voll automatischen Verifikation äußerst gering ist. Eine solche gewünschte *thematische Genauigkeit* ist in (BKG, 2009) definiert mit 95%, bezogen auf den Flächenanteil und die Anzahl der korrekten Objekte. Ziel ist es also, diese *thematische Genauigkeit* zu erreichen und gleichzeitig den manuellen Aufwand des menschlichen Bearbeiters zu minimieren. Daher wird in dieser Arbeit ein semi-automatisches Verifikationsverfahren benutzt. Alle GIS-Objekte, die vom System als richtig verifiziert werden, müssen dabei vom menschlichen Bearbeiter nicht mehr kontrolliert werden. Alle anderen GIS-Objekte werden vom menschlichen Bearbeiter überprüft. Die Entscheidung über Annahme oder Rückweisung eines vom System abgelehnten GIS-Objektes wird dann von dem menschlichen Bearbeiter durchgeführt. Wenn die geforderte *thematische Genauigkeit* nach dem Verifikationsprozess erreicht werden kann, ist eine höhere Anzahl von GIS-Objekten, die durch den menschlichen Betreuer begutachtet werden müssen, vertretbar, falls bei der Anwendung des Systems für den menschlichen Bearbeiter eine Zeitersparnis entsteht.

Schwerpunkt der Arbeit ist die Überprüfung der *thematischen Genauigkeit*. Die *geometrische Genauigkeit* der Grenzen der GIS-Objekte wird nicht überprüft und geht als Vorwissen in den Verifikationsprozess ein. Die *zeitliche Genauigkeit* des GIS wird, wie bereits beschrieben, implizit kontrolliert. Die Prüfung der *Konsistenz* als Maß der Qualität von Geodaten ist nicht Bestandteil der vorliegenden Arbeit.

Die Überprüfung der *thematischen Genauigkeit* von GIS-Datenbanken war bisher Fokus zahlreicher Arbeiten, u.a. (Eidenbenz et al., 2000), (Walter, 2004), (Busch et al., 2005). Bisherige Arbeiten mit dem Schwerpunkt der Verifikation konzentrierten sich jedoch nie auf GIS-Objekte der Objektklassen Acker- und Grünland. Der Fokus der Arbeiten, die sich bisher auf diese Objektklassen konzentrierten, lag in der Klassifikation, nicht aber in der Verifikation, u.a. (Ruiz et al., 2004), (Trias-Sanz, 2006), (Ruiz et al., 2007), (Balaguer et al., 2010) und (Recio et al., 2011). Die bisher veröffentlichten Arbeiten bezüglich der Klassifikation berücksichtigen nicht die Spezifikationen, die durch das GIS gegeben sind. Die hier vorgestellte Arbeit führt auf Grundlage einer Klassifikation die Verifikation von GIS-Objekten der Objektklassen Acker- und Grünland durch und berücksichtigt dabei Spezifikationen, die durch das GIS gegeben sind. So ist u.a. eine Segmentierung der Schläge innerhalb eines GIS-Objektes notwendig.

Bei der Klassifikation der Objektklassen Ackerland und/oder Grünland kommen spektrale, textuelle und/oder strukturelle Merkmale sowie geometrische- und weitere Merkmale zum Einsatz. Merkmale aus diesen Merkmalsgruppen werden oft in Mehrklassenproblemen genutzt, u.a. (Ruiz et al., 2004), (Trias-Sanz, 2006), (Ruiz et al., 2007), (Balaguer et al., 2010) und (Recio et al., 2011). Spezielle Verfahren zum Ableiten von Merkmalen, die insbesondere die Klassen Acker- und Grünland beschreiben, sind bisher nicht verwendet worden. Dies gilt besonders in Hinblick auf Verfahren zur Ableitung struktureller Merkmale.

Die meisten Klassifikationsansätze benutzen das Verfahren der größten Wahrscheinlichkeit (*Maximum-Likelihood-Verfahren*) oder baumförmige Klassifikationsverfahren (*Hierarchical Classifier* oder *Decision Tree*), speziell das C5-Klassifikationsverfahren nach Quinlan (2003). Es wird u.a. benutzt in (Ruiz et al., 2004), (Ruiz et al., 2007), (Balaguer et al., 2010) sowie (Recio et al., 2011). Das Verfahren der *Support Vector Machine* (Vapnik, 1995), das bereits in den unterschiedlichsten Bereichen der Wissenschaft und Technik Anwendung gefunden hat, u.a. in (Schölkopf et al., 1996), (Lu et al., 2001), (Nemmour und Chibani, 2006), (Fujimura et al., 2008), (Herbert und Riedel, 2010), (Römer und Plümer, 2010), hat für die Klassifikation von landwirtschaftlichen Nutzflächen bisher noch keine breite Anwendung gefunden. Dabei ist dieses Verfahren besonders gut für die Verifikation von GIS-Ackerland- und Grünlandobjekten geeignet, da u.a. die unterschiedlichen Erscheinungsformen von Ackerland berücksichtigt werden können. Auf Grund der unterschiedlichen Erscheinungsformen entstehen im Merkmalsraum räumlich getrennte Ballungen (Cluster), die mit der SVM klassifiziert werden können. Das Verfahren der SVM soll in dieser Arbeit zum Einsatz kommen.

1.4. Struktur der Arbeit

Im nächsten Kapitel wird zunächst ein Überblick über den Stand der Wissenschaft bezüglich der Schwerpunkte der Arbeit gegeben. Zum einen sollen Verfahren, die sich mit der Klassifikation landwirtschaftlicher Flächen beschäftigen, vorgestellt werden. Dabei wird sowohl auf die verwendeten Klassifikationsverfahren als auch -merkmale eingegangen. Zum anderen sollen Verfahren betrachtet werden, die sich mit der Verifikation von GIS beschäftigen. In diesem Teil wird ebenfalls eine Abgrenzung der Verifikation zum *update* und zur *change detection* erfolgen. Die meisten dieser Arbeiten beschäftigen sich jedoch nicht speziell mit landwirtschaftlichen Flächen. Als ein weiterer Punkt wird kurz auf die Segmentierung eingegangen, die benötigt wird, um die Schläge innerhalb eines GIS-Objektes zu bestimmen.

Im Anschluss daran wird im Kapitel 3 der in dieser Arbeit verwendete Klassifikations- und der Segmentierungsalgorithmus vorgestellt. Bei dem Klassifikationsverfahren handelt es sich wie bereits angeführt um das Verfahren der *Support Vector Machine* (SVM), deren Funktionsweise zunächst allgemein beleuchtet wird. Wichtige Erweiterungen, Anmerkungen zur praktischen Durchführung, die automatische Bestimmung der Parameter sowie Vergleiche mit anderen Klassifikationsverfahren sind ebenfalls Inhalt dieses Kapitels. Auf die Grundlagen des verwendeten Segmentierungsalgorithmus wird nur kurz eingegangen, da dieser auf Standardverfahren der Bildanalyse beruht.

Im Kapitel 4 wird dann die neue Methode zur semi-automatischen Verifikation von GIS-Objekten landwirtschaftlicher Objektklassen (speziell Acker- und Grünland) eingeführt. Dort wird neben dem Umgang

mit den GIS-Spezifikationen auf die zugrunde liegende Klassifikation und die dabei verwendeten Merkmale eingegangen.

Die Evaluierung der im Kapitel 4 vorgestellten Methode erfolgt im Kapitel 5. Nachdem die Kriterien zur Bewertung der Klassifikations- und Verifikationsergebnisse sowie die Datengrundlagen beschrieben wurden, werden zunächst die einzelnen Analyseschritte separat untersucht. Abschließend wird die Evaluation des gesamten Verfahrens durchgeführt.

Die Arbeit schließt mit einer Diskussion und einem Ausblick im Kapitel 6 ab.

2. Stand der Wissenschaft

Schwerpunkt der vorliegenden Arbeit ist die automatisierte Verifikation von GIS-Objekten der Objektarten Acker- und Grünland. Die meisten Arbeiten, die sich mit landwirtschaftlichen Nutzflächen beschäftigen, fokussieren nicht auf die Verifikation, sondern die Klassifikation, wobei die Klassifikation ein wesentlicher Bestandteil eines Verifikationsprozesses ist. Bei einer Klassifikation findet die Zuordnung von Pixeln oder Objekten zu vorher definierten Klassen statt (pixel- oder objektbasierte Klassifikation). Die Zuordnung zu einer Klasse erfolgt auf der Basis von pro Pixel/Objekt bestimmten Merkmale mittels eines gewählten Klassifikationsalgorithmus. Eine Klassifikation landwirtschaftlicher Nutzflächen ist komplex, da z.B. die landwirtschaftlichen Flächen Acker- und Grünland auf Grund unterschiedlicher natürlicher Gegebenheiten und landwirtschaftlicher Produktionsmethoden große regionale Unterschiede in ihrer Erscheinungsform aufweisen (Albertz, 2009). Die Anforderung an die Algorithmen besteht darin, diesen Umständen gerecht zu werden.

Da ein wesentlicher Teil des Verifikationsprozesses die Klassifikation der zu verifizierenden Klassen ist, stellt der erste Schwerpunkt dieses Kapitels einen Literaturüberblick zum Thema der Klassifikation landwirtschaftlicher Nutzflächen dar. Bevor auf die Merkmale für die Klassifikation landwirtschaftlicher Flächen eingegangen wird, werden zunächst die gängigen Klassifikationsverfahren, die für diese Aufgabe bisher verwendet wurden, vorgestellt. Es werden ebenfalls Vor- und Nachteile von pixel- und objektbasierter Klassifikation diskutiert. Anschließend werden die Merkmale, die zur Klassifikation verwendet werden, erörtert. Die Merkmale, die diskutiert werden, sind spektrale, textuelle, strukturelle oder geometrische Merkmale. Weitere Merkmale können ebenfalls für die Klassifikation verwendet werden. Solche Merkmale werden aus anderen Quellen wie einem Digitalen Geländemodell oder aus weiteren Erhebungen (z.B. weitere GIS-Datenbestände) abgeleitet (Smith und Fuller, 2001), (Recio et al., 2010), (Recio et al., 2011). Auf diese Merkmale kann in dieser Arbeit nicht zurückgegriffen werden, da sie nicht vorliegen. Aus diesem Grund wird auf die Gruppe der weiteren Merkmale nicht weiter eingegangen. Der Abschnitt schließt mit einer umfassenden Literaturübersicht zu Arbeiten unter Verwendung verschiedener Klassifikationsalgorithmen und verschiedener Merkmalsgruppen.

Im Abschnitt 2.2 wird auf andere Publikationen, die sich mit der Verifikation von GIS beschäftigen, eingegangen. Während bei der Klassifikation Spezifikationen des GIS nicht explizit berücksichtigt werden, geschieht dies bei der Verifikation. In diesem Abschnitt soll auch die Verifikation gegenüber den *updates* und der *change detection* abgegrenzt werden. Nachdem verschiedene Algorithmen zur Segmentierung der Schläge innerhalb eines GIS-Objektes in Abschnitt 2.3 diskutiert wurden, schließt das Kapitel mit einer Diskussion in Abschnitt 2.4.

In den in diesem Kapitel beschriebenen unterschiedlichen Veröffentlichungen werden allgemein gebräuchliche Qualitätsmaße wie *user's accuracy*, *producer's accuracy*, *overall accuracy* oder Kappa-Koeffizient verwendet (s. auch Abschnitt 5.1). Meist sind nicht alle Maße in den zitierten Veröffentlichungen angegeben. Die jeweils angegebenen Maße werden immer genannt.

2.1. Klassifikation von landwirtschaftlichen Nutzflächen

Die Voraussetzung für eine erfolgreiche Verifikation ist eine erfolgreiche Klassifikation der betrachteten GIS-Objekte. Während bei der Verifikation bisher landwirtschaftliche Nutzflächen kaum im Vordergrund stehen, gibt es zahlreiche Veröffentlichungen, die sich mit der Klassifikation solcher Flächen beschäftigen. Der Schwerpunkt des vorliegenden Literaturüberblicks liegt auf Methoden, die monotemporale multispektrale Bilddaten im Bereich von 0,5 bis 1 m Bodenauflösung verwenden, wie sie auch in dieser Arbeit im Mittelpunkt stehen. Der vorliegende Literaturüberblick ist nicht nur auf Acker- und Grünland beschränkt, sondern es wird auch auf Objektklassen eingegangen, die ähnliche regelmäßige Strukturen wie Ackerland enthalten. Eine Objektklasse, die dabei häufig im Mittelpunkt von Publikationen steht, sind Weinanbaugebiete.

In diesem Abschnitt sollen zunächst kurz die Klassifikationsverfahren und anschließend die Merkmale beschrieben werden, die bisher für diese Aufgabe eingesetzt wurden. Im Anschluss werden die bisherigen Veröffentlichungen zur Klassifikation landwirtschaftlicher Nutzflächen detailliert vorgestellt.

2.1.1. Klassifikationsalgorithmen

Es wird zwischen unüberwachten- und überwachten Klassifikationsverfahren unterschieden. In diesem Abschnitt werden Methoden beider Gruppen, die bereits für die Klassifikation landwirtschaftlicher Nutzflächen dienen, diskutiert. Weitere Methoden und Details sind in (Duda et al., 2000) und (Bishop, 2006) zu finden.

2.1.1.1. Unüberwachte Klassifikationsverfahren

Bei der unüberwachten Klassifikation werden Bildelemente mit ähnlichen Merkmalen zu Klassen zusammengefasst, ohne dass die Bedeutung der Klassen bekannt ist. Eine Klassenanzahl wird zwar vom Nutzer oft vorgegeben, die wahre Anzahl an Klassen ist jedoch unbekannt. Die Klassifikation läuft iterativ ab, bis ein Abbruchskriterium erreicht wird. Diese Verfahren werden bei der Klassifikation landwirtschaftlicher Nutzflächen, wenn überhaupt, als Vorverarbeitungsschritt eingesetzt. Sie spielen bei der Klassifikation landwirtschaftlicher Nutzflächen kaum eine Rolle.

Ein gängiges unüberwachtes Klassifikationsverfahren ist die Klassifikation mittels *k-means*. Bei diesem Verfahren werden zunächst K Cluster-Zentren im Merkmalsraum zufällig verteilt. Die Pixel/Objekte werden dem nächstgelegenen Cluster zugeordnet. Danach wird das Zentrum des Clusters neu ermittelt, z.B. durch Mittelbildung der zugehörigen Pixelpositionen. Die Zuordnung der Pixel/Objekte zu Cluster und die Berechnung der neuen Zentren der Cluster werden solange wiederholt, bis das Abbruchskriterium erreicht wird (Albertz, 2009). Die Nachteile dieses Verfahrens wurden bereits oben diskutiert. Das Verfahren wird als Vorverarbeitungsschritt in den Arbeiten von Fung und Chan (1994) verwendet.

2.1.1.2. Überwachte Klassifikationsverfahren

Bei der überwachten Klassifikation sind die Anzahl der Klassen und deren Bedeutung bekannt. Außerdem sind Trainingsgebiete gegeben. Für die Trainingsgebiete ist die Klassenzugehörigkeit von Pixeln/Objekten bekannt. Der Klassifikationsalgorithmus lernt mittels der Trainingsdaten die Lage und Ausdehnung der Cluster im Merkmalsraum und kann neue Pixel/Objekte, deren Merkmale bekannt sind, dann einer Klasse zuordnen.

Ein gängiges Verfahren in der Fernerkundung ist das Verfahren der größten Wahrscheinlichkeit (*Maximum-Likelihood-Verfahren*, kurz ML-Verfahren). Das ML-Verfahren berechnet auf Grund statistischer Kenngrößen (z.B. Erwartungswert oder Standardabweichung) der vorgegebenen Klassen die Wahrscheinlichkeit $P(\mathbf{x}|C)$ mit der jedes Pixel/Objekt auf Grund dessen Merkmalen $\mathbf{x}$ der Klasse C angehört. Ein Pixel/Objekt wird der Klasse mit der größten Wahrscheinlichkeit $\max. P(\mathbf{x}|C)$ zugewiesen. Dabei geht man davon aus, dass die Daten jeder Klasse einer bestimmten Verteilung (z.B. der Normalverteilung) unterliegen (Albertz, 2009). Wird diese Verteilungsannahme verletzt, kann es sein, dass der ML Klassifikator inkonsistent ist. Die Parameter der Wahrscheinlichkeitsdichtefunktion werden aus den Trainingsdaten ermittelt. Bei dem ML-Verfahren werden alle Pixel/Objekte einer Klasse zugeteilt. $\max. P(\mathbf{x}|C)$ entspricht dem $\max. P(C|\mathbf{x})$ des Bayestheorems, wenn beim Bayestheorem $P(C|\mathbf{x})$ mit $P(C|\mathbf{x}) = P(\mathbf{x}|C)P(C)/P(\mathbf{x})$ die Wahrscheinlichkeit des Auftretens einer Klasse $P(C)$ unbekannt ist und in diesem Fall $P(C)$ für alle Klassen gleich angenommen wird. Das ML-Verfahren führt i.d.R. zu guten Ergebnissen (Albertz, 2009), kann jedoch mit räumlich getrennten Clustern im Merkmalsraum, die einer Objektklasse angehören, nicht umgehen. Gerade dies ist bei den unterschiedlichen Erscheinungsformen von Ackerland jedoch notwendig. Das ML-Verfahren wird für die Klassifikation landwirtschaftlicher Flächen u.a. in (Fung und Chan, 1994), (Warner und Steinmaus, 2005), (Itzerott und Kaden, 2006) und (Itzerott und Kaden, 2007) verwendet.

Ein Spezialfall des ML-Verfahrens ist das *Linear Discriminant Function Verfahren* (kurz LDF-Verfahren), bei dem die Normalverteilung und gleiche Kovarianzmatrizen vorliegen. Die genannten Nachteile des ML-Verfahrens treffen auch auf das LDF-Verfahren zu, weil dieses im Wesentlichen auf dem ML-Verfahren basiert. Das LDF-Verfahren findet Anwendung bei der Klassifikation unterschiedlicher landwirtschaftlicher Nutzflächen in (Haralick et al., 1973).

Bei dem Verfahren der nächsten Nachbarschaft (*Minimum-Distance-Verfahren*, kurz MD-Verfahren) werden zunächst für die Trainingsgebiete jeder Klasse die Mittel aller Merkmale berechnet. Für jedes zu klassifizierende Pixel/Objekt wird anschließend der Abstand im Merkmalsraum zu den Mittelpunkten aller Cluster berechnet. Das Pixel/Objekt wird der Klasse zugeordnet, zu deren Mittelpunkt der Abstand am kürzesten ist. Als Distanzmaß kann nicht nur die euklidische Distanz, sondern auch jedes andere Distanzmaß verwendet werden. Unterschiedliche Streubereiche der Klassen werden nicht berücksichtigt (Ausnahme:

Mahalanobisdistanz). Deshalb kann ein Pixel/Objekt einer Klasse zugewiesen werden, der es mit großer Wahrscheinlichkeit nicht zugehört (Albertz, 2009).

Beim Nächsten-Nachbarn-Verfahren (*k-Nearest-Neighbor-Verfahren*, kurz kNN-Verfahren) ist nicht der geringste Abstand zu dem Mittel, wie beim MD-Verfahren, Grund für die Zuordnung zu einer Klasse, sondern ein Pixel/Objekt wird der Klasse zugeordnet, die am häufigsten in seiner k nächsten Nachbarschaft auftritt. Nachteil ist, ähnlich wie beim ML-Verfahren, dass mit räumlich getrennten Clustern im Merkmalsraum, die derselben Klasse angehören, nicht umgegangen werden kann.

Bei Box-Klassifikator (*Box Classifier*), oder auch Parallelepiped-Klassifikator genannt, werden zunächst die minimalen und maximalen Merkmalswerte der Klassen im Merkmalsraum bestimmt und diese dann genutzt, um im Merkmalsraum Quader zu definieren, siehe u.a. (Albertz, 2009). Liegt ein Merkmal eines Pixels/Objektes mit unbekannter Klassenzugehörigkeit innerhalb des Quaders, wird dieses dieser Klasse zugeordnet. Ein Vorteil des Box-Klassifikators liegt darin, dass dieser sehr schnell ist. Pixel/Objekte können jedoch nicht immer einer Klasse zugewiesen werden, nämlich dann, wenn das Merkmal in einer Quaderlücke liegt. Es kann also auch nicht klassifizierte Pixel/Objekte geben. Des Weiteren sind Überlagerungsbereiche der Quader möglich. Betroffene Pixel/Objekte, die im Überlagerungsbereich liegen, können nicht eindeutig einer Klasse zugewiesen werden. Ein Box-Klassifikator kommt in den Arbeiten von Itzerott und Kaden (2006, 2007) sowie von Wassenaar et al. (2002) zum Einsatz.

Bei der hierarchischen oder baumförmigen Klassifikation (*Hierarchical Classifier* oder *Decision Tree*) erfolgt die Zuweisung eines Pixels/Objektes nicht in einem einmaligen Vorgang, sondern schrittweise durch eine Folge von Einzelentscheidungen. Auf jeder Stufe findet eine Entscheidung zwischen nur wenigen Möglichkeiten (oft nur zwei) statt. Die Entscheidung zwischen verschiedenen Möglichkeiten erfolgt solange, bis eine endgültige Klasse erreicht ist, u.a. (Albertz, 2009). Der Baum wird mittels Trainingsdaten erstellt. Das Verfahren ist sehr flexibel anwendbar, da für jede Einzelentscheidung die günstigste Kombination und das zweckmäßigste Kriterium benutzt werden kann. Es neigt dadurch aber zur Überanpassung an die Trainingsdaten. Eine Art eines Decision Tree Verfahrens, das in der Klassifikation landwirtschaftlicher Nutzflächen Anwendung findet, ist das C5-Verfahren (Quinlan, 2003). Es wird u.a. benutzt in (Ruiz et al., 2004), (Ruiz et al., 2007), (Balaguer et al., 2010) und (Recio et al., 2011).

Einen anderen Klassifikationsalgorithmus stellt die *Support Vector Machine* (SVM; Vapnik, 1998) dar. Die SVM ist ein binäres Klassifikationsverfahren, das zwei Klassen an Hand ihrer Merkmale durch eine Hyperebene im Merkmalsraum so trennt, dass der Abstand derjenigen Merkmalsvektoren, die der Hyperebene am nächsten liegen, maximiert wird. Liegen nicht linear trennbare Daten im Merkmalsraum vor, wird durch eine Transformation der Daten in einen höher dimensional Merkmalsraum eine lineare Trennbarkeit erreicht. Die Transformation muss nicht explizit durchgeführt werden, sondern wird durch den sogenannten Kernel-Trick erreicht. Die SVM kann auch auf Mehrklassenprobleme angewendet werden, indem verschiedene Zwei-Klassen-SVM miteinander kombiniert werden, was auf Kosten der Laufzeit geschieht. Die SVM wurde für die Trennung landwirtschaftlicher Nutzflächen bisher nicht verwendet, kommt aber in dieser Arbeit zum Einsatz. Daher wird auf die Funktionsweise sowie Vor- und Nachteile der SVM in Kapitel 3 näher eingegangen.

Ein weiteres Verfahren, das für die Klassifikation von Luft- und Satellitenbildern angewendet wird, ist das Verfahren der Markov-Zufallsfelder (*Markov Random Fields*, kurz MRF). Bei den MRF wird das Bild als ungerichteter Graph interpretiert, bei dem die Knoten die Pixel/Objekte und die Kanten Beziehungen zu Nachbarpixeln/-objekten repräsentieren. Auf Grund der Berücksichtigung der Klassenzugehörigkeit der Nachbarpixel/-objekte können lokale Zusammenhänge berücksichtigt werden. Das zu klassifizierende Pixel/Objekt wird damit nicht nur auf Grund seiner Merkmale, sondern auch auf Grund der Klassenzugehörigkeit seiner Nachbarpixel/-objekte einer Klasse zugeordnet. Somit kann dem bei bisher genannten Klassifikationsverfahren oft auftretenden „Salt-and-Pepper“ Effekt entgegengetreten werden. Nachteile sind, dass die Interaktionen auf die Klassen der benachbarten Pixel/Objekte beschränkt sind und dass es sich um ein sehr rechenintensives Verfahren handelt. Detaillierte Informationen können (Li, 2009) entnommen werden. Ein MRF für die Klassifikation landwirtschaftlicher Flächen wird u.a. in der Arbeit von Busch et al. (2005) verwendet.

Während bei den Klassifikationsverfahren mit MRF nur die Information der Klassenzugehörigkeit der lokalen Nachbarn zur Bestimmung der Klasse eines Pixels/Objektes Berücksichtigung finden, fließen bei den *Conditional Random Fields* (CRF; Kumar und Hebert, 2006) auch die Merkmale der Nachbarn in die

Klassifikationsentscheidung mit ein. Einem Nachteil der MRF konnte bei der CRF somit entgegengetreten werden, das geht jedoch auf Kosten der Laufzeit. Die CRF werden für die Klassifikation u.a. in (Hoberg und Rottensteiner, 2010) angewendet.

2.1.1.3. Pixel- und objektbasierte Klassifikation

Eine Klassifikation kann generell pixel- oder objektbasiert durchgeführt werden. Pixelbasierte Klassifikationsansätze führen eine Klassifikation auf Grund von Merkmalen eines Pixels und/oder Merkmalen von dessen lokaler Umgebung durch. Eine pixelbasierte Klassifikation wird z.B. in (Alpin et al., 1999), (Hirose et al., 2004), (Delenne et al., 2007) und (Ranchin et al., 2001) durchgeführt. Ist das Ziel wie in dieser Arbeit die Verifikation eines GIS-Objektes, muss nach der pixelbasierten Klassifikation deren Ergebnis auf das GIS-Objekt übertragen werden. Diese Art der Klassifikation von GIS-Objekten wird auch „post-object-classification“ (Janssen und van Amsterdam, 1991) genannt. Aus der Überführung des pixelbasierten Klassifikationsergebnisses auf ein GIS-Objekt können sich Nachteile ergeben. Wenn z.B. ein GIS-Objekt aus verschiedenen klassifizierten Pixeln besteht, kommt es bei der Überführung des pixelbasierten Klassifikationsergebnisses auf das GIS-Objekt zu Generalisierungsproblemen. Besondere Schwierigkeiten treten auf, wenn das pixelbasierte Klassifikationsergebnis starke Eigenschaften eines „Salt-And-Pepper-Effect“ besitzt (Blaschke, 2010). Dieser Effekt kann durch die Verwendung von MRF verringert werden. Durch das Auftreten verschieden klassifizierter Pixel innerhalb eines GIS-Objektes können aber auch Nutzungsarten, die innerhalb eines GIS-Objektes im Sinne der Verifikation unter Berücksichtigung von GIS-Spezifikation nicht erlaubt sind, detektiert werden.

Bei der objektbasierten Klassifikation wird die Klassifikation auf Grund von Merkmalen durchgeführt, die für ein Objekt bzw. Segment² abgeleitet werden, wobei ein GIS-Objekt aus mehreren Segmenten bestehen kann. Eine Klassifizierung, die mittels direkt von einem Segment abgeleiteten Merkmalen durchgeführt wird, wird auch „pre-object-classification“ genannt (Janssen und van Amsterdam, 1991). Durch die Verwendung von Merkmalen, die sich auf ein Segment beziehen, wird der „Salt-And-Pepper-Effect“ belanglos (Blaschke, 2010). Nach Löcherbach (1994) sowie Alpin und Smith (2008) ist das Ergebnis der objektbasierten Analyse näher an dem eigentlichen Ziel, der Klassifikation eines GIS-Objektes, als die pixelbasierte Klassifikation. Falls ein GIS-Objekt aus mehreren Teilsegmenten besteht, müssen die Teilergebnisse der Segmente für das Klassifikationsergebnis des GIS-Objektes fusioniert werden. Ein Vorteil der objektbasierten Klassifikation ist, dass Merkmale, die das Segment näher beschreiben, u.a. die Form, als Merkmale in die Klassifikation eingeführt werden können. Ein weiterer Vorteil ist, dass mehrere Pixel zur Bestimmung der Merkmale beitragen und das Ergebnis somit robuster ist. Beispiele für objektbasierte Klassifikationen sind in (Smith und Fuller, 2001) und (Müller et al., 2010) gegeben.

Nachteile der objektbasierten Klassifikation bestehen, wenn eine Untersegmentierung vorliegt, also innerhalb eines Segmentes mehrere Objektklassen auftreten (Alpin et al., 1999). In diesem Fall bilden sich die Klassengrenzen nicht als Segmentgrenzen ab. Ein Beispiel ist, dass in einem Ackerland-segment, das an eine Siedlung grenzt, ein Haus neu entstanden ist und dieses im Ackerlandsegment liegt und kein eigenes Segment bildet. Nimmt das Haus nur eine geringe Fläche gegenüber dem Reststück des Segmentes ein, unterscheiden sich die Merkmale des Segmentes mit und ohne Haus nicht signifikant voneinander. Wenn Häuser in einem Ackerlandsegment nicht erlaubt sind, würde dies zu einem falschen Klassifikationsergebnis führen. Ein weiteres Beispiel für eine Untersegmentierung liegt vor, wenn ein Segment aus mehreren Teilsegmenten besteht, die sich in ihren Merkmalen unterscheiden, aber derselben Klasse angehören (z.B. Mais und Weizen der Objektklasse Ackerland). Die Merkmalswerte des gesamten Segmentes könnten mit den Merkmalswerten eines Segmentes einer anderen Objektklasse (z.B. Grünland) übereinstimmen. Auch in einem solchen Fall wird das Segment falsch klassifiziert. Im Gegensatz zur objektbasierten Klassifikation können bei einer pixelbasierten Klassifikation Fehler, bei der eine Teil eines Segmentes eine andere Nutzung beinhaltet, z.B. ein Haus, entdeckt werden. Jedoch ist die Unterscheidung zwischen dem unerwünschten „Salt-And-Pepper-Effect“ und einer tatsächlichen korrekten Detektion einer kleinen Teilfläche mit einer anderen Nutzung innerhalb eines Segmentes bei der notwendigen Generalisierung für die pixelbasierte Klassifikation eine Herausforderung.

² In der Literatur wird der Begriff *Objekt* oft für ein Bildsegment benutzt, dessen Pixel ähnliche Eigenschaften besitzen. Die Grenze eines solchen Segmentes stimmt in den meisten Fällen nicht mit der Grenze eines GIS-Objektes überein. Es muss daher zwischen einem Segment und einem GIS-Objekt unterschieden werden. Wenn in dieser Arbeit ein GIS-Objekt gemeint ist, wird dies explizit angegeben. Wird der Begriff Segment (oder Objekt) benutzt, ist ein Segment gemeint, dessen Pixel ähnliche Eigenschaften besitzt.

Eine Kombination beider Verfahren (pixel- und objektbasiert) ist daher sinnvoll, um die Vorteile beider Verfahren zu nutzen.

2.1.2. Merkmale für die Klassifikation

Ein Überblick über Klassifikationsmethoden ist in (Lu und Weng, 2007) gegeben. Lu und Weng (2007) betonen die Herausforderung der Klassifikation fernkundlicher Daten, die von vielen Faktoren beeinflusst werden, wie z.B. die Komplexität der Landschaft. Ebenfalls wird herausgestellt, dass neben den spektralen Merkmalen textuelle Merkmale und Kontextinformationen mit steigender räumlicher Auflösung an Bedeutung gewinnen. Bei der Klassifikation landwirtschaftlicher Flächen kommen zusätzlich auch textuelle, strukturelle und geometrische Merkmale sowie weitere Merkmale zum Einsatz.

2.1.2.1. Spektrale Merkmale

Spektrale Merkmale enthalten Informationen über die Reflexionseigenschaften von Objekten und hängen nicht nur von der Objektklasse selbst, sondern auch von anderen Bedingungen wie dem Stadium der Vegetation (phänologische Änderungen), dem unterschiedlichen Feuchtigkeitsgehalt der Blätter sowie dem Stand der Sonne oder der Bodenbeschaffenheit zum Zeitpunkt der Aufnahme ab. Spektrale Merkmale basieren auf der vom Sensor aufgezeichneten elektromagnetischen Strahlung (Albertz, 2009).

Man unterscheidet einzelne Systeme der Datenaufnahme nach den Wellenbereichen der aufgezeichneten elektromagnetischen Strahlung, wobei die betreffenden Spektralbereiche Kanäle genannt werden. Von einem multispektralen System spricht man, wenn gleichzeitig mehrere Messwerte in verschiedenen Wellenlängenbereichen erfasst und aufgezeichnet werden (Albertz, 2009). Zum Beispiel zeichnet IKONOS multispektrale Bilder mit vier Kanälen auf, deren Wellenlängenbereiche in Tabelle 1 zusammengefasst sind.

Kanal	Wellenlänge [µm]
(1) Blau	0,445- 0,516
(2) Grün	0,506- 0,595
(3) Rot	0,632- 0,698
(4) Nahes Infrarot	0,757- 0,853

Tabelle 1: Wellenlängenbereiche eines multispektralen IKONOS-Bildes.

An Hand unterschiedlicher spektraler Merkmale ist die Trennung von Objekten verschiedener Klassen möglich. Die spektralen Merkmale können direkt oder indirekt zur Klassifizierung verwendet werden. Bei einer direkten Verwendung der Informationen wird direkt auf die vom Sensor registrierten Daten zurückgegriffen. Eine Möglichkeit, die Reflexionseigenschaften indirekt zur Trennung unterschiedlicher Objektklassen zu benutzen ist die Ableitung von Merkmalen aus den Farbkanälen, wie z.B. des *Normalized Difference Vegetation Index* (NDVI), der sich besonders gut für die Trennung von Flächen mit und ohne Vegetation eignet (Lillesand und Kiefer, 2000).

Häufig verwendete spektrale Merkmale pro Segment bzw. innerhalb der lokalen Nachbarschaft eines Pixels werden jeweils für jeden Farbkanal und wenn möglich auch für den NDVI bzw. andere Indices berechnet und können sein:

- Mittelwert,
- Standardabweichung,
- Maximum und Minimum,
- Differenz aus Maximum und Minimum,
- Am häufigsten auftretender Wert,
- Medianwert.

Wenn Segmente aus einer Klasse bestehen, sind die berechneten spektralen Merkmale repräsentativ für die Klasse innerhalb des Segmentes. Betrachtet man nur die lokale Umgebung eines Pixels, so können verschiedene Klassen in der Umgebung auftreten. Eine Zuweisung des zentralen Pixels zu einer Klasse kann daher zu einer Falschklassifizierung des Pixels führen.

Spektrale Merkmale werden in zahlreichen Veröffentlichungen verwendet, darunter auch in den Arbeiten von Haralick et al. (1974), Fung und Chan (1994), Ruiz et al. (2004, 2007), Itzerott und Kaden (2006, 2007), Marçal und Cunha (2007), Balaguer et al. (2010), Peled und Gilichinsky (2010) sowie Recio et al. (2011). Sie bieten ein gutes Mittel zur Klassifikation landwirtschaftlicher Nutzflächen. Bezogen auf Acker- und Grünland

gibt es jedoch innerhalb eines Jahres Zeitpunkte, bei denen sich die spektralen Merkmale dieser beiden Klassen nicht unterscheiden. Die alleinige Verwendung spektraler Merkmale für die Klassifikation von Acker- und Grünland ist also insbesondere für monotemporale Aufnahmen problematisch.

2.1.2.2. Texturelle Merkmale

Das Wort „Textur“ stammt vom lateinischen Wort „texere“ bzw. „textura“ ab und bedeutet so viel wie „Gewebe“. Die Definition nach Tönnies (2005, S. 213) lautet: „Unter der Textur einer Region versteht man eine gemeinsame Eigenschaft der Grauwertverteilung dieser Region, die auf eine Bildungsregel zurückzuführen ist.“ Die unterschiedlichen Variationen der Grauwerte, d.h. die Textur bzw. die enthaltenen Muster unterscheiden sich für verschiedene Objektklassen in Luft- und Satellitenbildern. Mit Hilfe der Textur, oder besser Merkmalen, die die Textur beschreiben, ist es möglich, zwischen unterschiedlichen Objektklassen zu differenzieren.

Häufige Methoden, die zur Beschreibung der Textur einer Region eingesetzt und zur Ableitung von Merkmalen für die Klassifikation landwirtschaftlicher Nutzflächen verwendet werden, sind:

- *Grey Level Co-Occurrence Matrix* (GLCM)
- *Neighborhood Gray-Tone Difference Matrix* (NGTDM)

Die GLCM wurde von Haralick et al. (1973) eingeführt. Die GLCM beschreibt die Häufigkeit des gemeinsamen Vorkommens von Grauwerten von Pixeln mit Abstand Δ in Richtung α . Aus der GLCM können Maße bestimmt werden, die die Textur einer Region näher beschreiben. In Haralick et al. (1973) werden insgesamt 14 Texturmaße vorgestellt, von denen die am häufigsten verwendeten Maße Energie, Kontrast, Homogenität und Korrelation sind. Detaillierte Informationen über das Aufstellen der GLCM sowie das Ableiten von Maßen aus der GLCM zur Klassifikation können (Haralick et al., 1973) und (Tönnies, 2005) entnommen werden. Die Haralick-Merkmale werden u.a. in der Veröffentlichung von Haralick et al. (1973) selbst sowie in den Veröffentlichungen von Smits und Annoni (1999), Gong et al. (2003), Ruiz et al. (2004, 2007), Balaguer et al. (2010), Müller et al. (2010) und Recio et al. (2011) zur Klassifikation von landwirtschaftlichen Flächen benutzt. Dabei konnten immer nur landwirtschaftliche Flächen von nicht landwirtschaftlichen Flächen abgegrenzt werden. Die Unterscheidung zwischen verschiedenen landwirtschaftlichen Flächen mit Hilfe von Merkmalen, die ausschließlich aus der GLCM abgeleitet wurden, konnte keine zufriedenstellenden Ergebnisse erbringen.

Das Verfahren der NGTDM wurde von Amadasun und King (1989) eingeführt. Bei der NGTDM spielen die Differenzen der Grauwerte der jeweiligen Bildpunkte und der Grauwerte der jeweiligen lokalen Nachbarn eine wichtige Rolle. Die NGTDM ist definiert als ein Vektor und hat so viele Einträge wie Grauwerte im Bild vorhanden sind. Der i -te Eintrag der NGTDM wird als die Summe der Grauwertdifferenzen zwischen allen Bildpunkten mit dem Grauwert i und deren jeweiligen lokalen Nachbarn bestimmt. Aus der NGTDM lassen sich ähnlich wie bei der GLCM verschiedene Texturmaße wie z.B. Komplexität und Texturstärke ableiten. Detaillierte Informationen zur Aufstellung der NGTDM sind (Amadasun und King, 1989) zu entnehmen. Die NGTDM und daraus abgeleitete Texturparameter wurden in der Arbeit von Rengers und Prinz (2009) für die Klassifikation von Fernerkundungsdaten genutzt. Auch hier konnte ähnlich wie bei der GLCM zwar zwischen landwirtschaftlich und nicht-landwirtschaftlich genutzten Flächen unterschieden werden, zwischen unterschiedlich landwirtschaftlich genutzten Flächen (hier: Acker- und Grünland) konnte jedoch nicht getrennt werden.

2.1.2.3. Strukturelle Merkmale

Strukturelle Merkmale sind Merkmale, die spezielle Strukturen in einer Region (z.B. Bearbeitungsstrukturen) beschreiben. Basierend auf diesen Merkmalen ist die Klassifikation der Region möglich. Strukturelle Merkmale können mittels unterschiedlicher Techniken ermittelt werden, z.B.

- Semi-Variogramm,
- Autokorrelation,
- Houghtransformation,
- Fouriertransformation,
- Gradientenhistogramme.

Das Semi-Variogramm wird verwendet, um die räumliche Variation von Grauwerten in Bildern zu messen. Es ist ein Diagramm, in dem die Semivarianz und Distanz gegeneinander aufgetragen werden. Nähere Informationen, eine allgemeine Herleitung sowie eine Einführung zur Nutzung von Semi-Variogrammen in

der Fernerkundung sind in (Wackernagel, 2003), (Woodcock et al., 1988a) und (Woodcock et al., 1988b) sowie in (Curran, 1988) und (Curran, 2001) gegeben. Für die Extraktion von Merkmalen zur Klassifikation landwirtschaftlicher Flächen werden in (Ruiz et al., 2004), (Trias-Sanz, 2006), (Ruiz et al., 2007), (Balaguer et al., 2010) und (Recio et al., 2011) Merkmale aus dem Semi-Variogramm abgeleitet und diese dann erfolgreich für die Klassifikation eingesetzt. Trias-Sanz (2006) benutzt ausschließlich Merkmale, die aus einem Semi-Variogramm für die Trennung unterschiedlicher landwirtschaftlicher Flächen abgeleitet sind und kann zufriedenstellende Ergebnisse damit erzielen. Diese Arbeit zeigt aber auch, dass es zu Problemen kommen kann, wenn in einem Segment eine nicht homogene Fläche vorliegt, z.B. wenn sich der Abstand der Strukturen im Segment oder die Hauptbewirtschaftungsrichtung selbst ändert.

Die Autokorrelation beschreibt die Korrelation eines Signales mit sich selbst als Funktion des Abstandes, d.h. wie stark ein Signal mit sich selbst korreliert ist. Die Autokorrelation wird meist auf zeitliche Signale angewandt, kann aber auch auf räumliche Signale wie Bilder angewandt werden (Cliff und Ord, 1981). Starke Autokorrelation ist in Bildern mit regelmäßigen Strukturen zu finden, wie z.B. einer Plantage. Das Verfahren mittels Autokorrelation ist besonders in Gebieten geeignet, die regelmäßig wiederkehrende Strukturen besitzen, es reagiert jedoch sehr sensibel auf Ungleichmäßigkeiten wie z.B. Abweichungen der Abstände zwischen Bäumen einer Plantage oder Bewirtschaftungsspuren eines Ackerlandes. Wenn z.B. Bäume in einer Plantage oder Bewirtschaftungsspuren im Bild eines Ackerlandes nicht zu erkennen sind, kann mit dem Verfahren kein zuverlässiges Ergebnis erreicht werden. Ein Beispiel für die Anwendung der Autokorrelation zur Klassifikation von landwirtschaftlichen Flächen, speziell Weinanbauflächen und Plantagen in IKONOS-Bildern, wird in (Warner und Steinmaus, 2005) vorgestellt.

Die Houghtransformation ist eine Technik zur Detektion von parametrisierbaren Objekten wie Linien und Kreisen aus Bildern unter Verwendung eines geeigneten Akkumulationsraumes (genannt Hough-Raum). Als Eingangsbild dient ein binäres Bild, das die Objekte von Interesse enthält (z.B. ein Kantenbild mit Bearbeitungsspuren eines Ackerlandobjektes). Im Hough-Raum sind die parametrisierten Objekte als Maxima zu erkennen. Zum Beispiel können Linien mittels Winkel zwischen Normalenvektor und der x-Achse (φ) sowie dem Abstand der Linie vom Koordinatenursprung (d) parametrisiert werden. Erscheinen also im Hough-Raum mehrere Maxima mit demselben φ , dann liegen im Eingangsbild mehrere parallele Linien vor. Dabei können die Linien unterbrochen sein. Nähere Informationen sind u.a. (Jähne, 2002), (Tönnies, 2005) zu entnehmen. Ist das Eingabebild ein binäres Kantenbild einer Ackerfläche, so können auf diese Weise die Hauptbewirtschaftungsrichtung der Ackerfläche und auch die Abstände zwischen den Bewirtschaftungsspuren bestimmt werden. Aus dem Hough-Raum werden dann Merkmale abgeleitet (z.B. nicht nur das Vorhandensein einer Hauptbewirtschaftungsrichtung, sondern auch die mittleren Abstände zwischen den Bewirtschaftungsspuren), die zur Klassifikation von landwirtschaftlichen Nutzflächen dienen, was u.a. in den Arbeiten von Ruiz et al. (2004, 2007), Balaguer et al. (2010) und Recio et al. (2011) dargestellt wird. Merkmale, abgeleitet aus einem Spezialfall der Houghtransformation, werden in (Chanussot et al., 2005) und (Trias-Sanz, 2006) zur Klassifikation landwirtschaftlicher Flächen genutzt. Die Houghtransformation setzt immer die erfolgreiche Detektion der Objekte von Interesse (z.B. der Bearbeitungsspuren) voraus, wobei die Technik nicht sensibel auf z.B. kleine Lücken im Kantenbild reagiert. Die Ableitung der Merkmale aus dem Hough-Raum stellt sich dabei als kompliziert dar, was die Zuverlässigkeit der abgeleiteten Merkmale beeinflussen kann (Trias-Sanz, 2006).

Die Fouriertransformation zerlegt ein Signal wie ein Bild in eine Summe von Sinus- und Cosinus-Kurven unterschiedlicher Amplitude, Phase und Frequenz. Nähere Angaben sind u.a. in (Tönnies, 2005) gegeben. Sind in einem Bild regelmäßige Strukturen wie Bearbeitungsspuren in Ackerflächen oder Baumreihen in Plantagen vorhanden, so ist dies im Amplitudenbild an Hand von Maxima zu erkennen. Aus dem Amplitudenbild können daher Merkmale abgeleitet werden, durch die eine Unterscheidung von Objektklassen möglich ist. Fouriertransformationen zur Detektion von regelmäßigen Strukturen werden in (Wassenaar et al., 2002), (Chanussot et al., 2005) und (Delenne et al., 2007) angewendet. Spezielle Formen der Fouriertransformation werden in (Rabatel et al., 2008), (Delenne et al., 2008), und in (Ranchin et al., 2001) für die Klassifikation landwirtschaftlicher Flächen verwendet. Ähnlich wie bei dem Verfahren der Autokorrelation sind auch Verfahren, basierend auf der Fouriertransformation, beschränkt auf Objekte, die streng gleichmäßige Strukturen enthalten, was bei Ackerland nicht immer gegeben ist. Für die Ableitung von zuverlässigen Merkmalen aus dem Amplitudenbild muss oft Vorwissen, wie der Abstand zwischen Pflanzreihen, bekannt sein, um die Suche nach Maxima einschränken zu können. Dies ist bei manchen Objektarten wie Weinanbauflächen oft gegeben. Bei Ackerland hingegen können die Abstände der

Pflanzreihen bzw. Bewirtschaftungsspuren je nach Fruchtart und verwendeter Bewirtschaftungsmaschinen stark schwanken und stehen im Normalfall nicht zur Verfügung.

Das Gradientenhistogramm (*histogram of gradients*) stellt die Häufigkeit des Vorkommens einer Gradientenrichtung (*gradient orientation*), gewichtet durch die Gradientenstärke (*gradient magnitude*), dar. Grundlage für die Berechnung des Gradientenhistogramms ist eine vorausgegangene Berechnung der Gradientenrichtung und –stärke. Besitzt ein Bild viele starke Kanten mit derselben Richtung, bildet sich im Gradientenhistogramm ein Maximum heraus. Merkmale, die aus dem Gradientenhistogramm abgeleitet werden, können in die Klassifikation landwirtschaftlicher Flächen eingehen; Beispiele sind in (Ruiz et al., 2004), (Ruiz et al., 2007), (Balaguer et al., 2010), (Hoberg und Rottensteiner, 2010) und (Recio et al., 2011) zu finden. In den Publikationen von Ruiz et al. (2004, 2007), Balaguer et al. (2010) und Recio et al. (2011) werden genauso wie in (Hoberg und Rottensteiner, 2010) Merkmale basierend auf Gradientenhistogrammen neben Merkmalen aus weiteren Verfahren genutzt. Es ist jedoch anzumerken, dass der Schwerpunkt der Arbeit von Hoberg und Rottensteiner (2010) auf der Trennung von *Siedlung* und *Nicht-Siedlung* lag. In (Hoberg und Müller, 2011) werden landwirtschaftliche Flächen klassifiziert. Hoberg und Müller (2011) verzichten gegenüber Hoberg und Rottensteiner (2010) auf Merkmalen aus dem Gradientenhistogramm, da sie nach eigenen Angaben keine Verbesserung der Ergebnisse mit diesen Merkmalen für landwirtschaftliche Flächen erzielen konnten. Stattdessen werden in (Hoberg und Müller, 2011) spektrale Merkmale aus multispektralen Bilddaten verwendet.

2.1.2.4. Geometrische Merkmale

Weitere Merkmale können aus den vorliegenden GIS-Daten abgeleitet werden. Lu und Weng (2007) stellen heraus, dass die Integration von Fernerkundungsdaten und GIS-Daten signifikant zur Klassifikationsverbesserung beiträgt. Auch Wilkinson (1996) beschreibt die Nutzung von GIS-Daten als ergänzende Information, um abgeleitete Produkte aus Fernerkundungsdaten zu verbessern.

Merkmale, die aus dem GIS abgeleitet werden können, sind u.a. folgende Merkmale eines GIS-Objektes:

- Form,
- Kompaktheit,
- Fläche,
- Umfang.

In den Arbeiten von Ruiz et al. (2004, 2007), Balaguer et al. (2010) sowie Recio et al. (2011) werden Formmerkmale verwendet, um das Ergebnis einer Klassifikation unterschiedlicher landwirtschaftlicher Flächen zu verbessern. Eine Klassifikation landwirtschaftlicher Nutzflächen an Hand reiner Formmerkmale wurde noch nicht durchgeführt. Der Erfolg eines solchen Ansatzes ist zu bezweifeln, da die Form nicht immer die Eigenschaften eines GIS-Objektes zuverlässig beschreibt.

Formmerkmale können nicht nur aus dem GIS abgeleitet, sondern auch für Segmente, die mittels Segmentierung bestimmt wurden, berechnet werden, wie dies z.B. in der Arbeit von Peled und Gilichinsky (2010) geschieht.

2.1.3. Beispiele für die Klassifikation landwirtschaftlicher Flächen

Nach den verschiedenen Klassifikationsansätzen sowie Merkmalen, die zur Klassifikation landwirtschaftlicher Flächen eingesetzt werden, können nun die Ansätze zur Klassifikation landwirtschaftlicher Nutzflächen detailliert beschrieben werden. Die Ansätze sind gruppiert nach den Merkmalen, die für die Klassifikation verwendet werden, beginnend mit Verfahren, die Merkmale aus nur einer, anschließend zweier und abschließend mehr als zwei Merkmalsgruppen benutzen.

2.1.3.1. Klassifikation mittels spektraler Merkmale

Rein spektrale Merkmale für die Klassifikation landwirtschaftlicher Flächen kommen unter Verwendung von monotemporalen multispektralen hochaufgelösten Bilddaten mit einer geometrischen Auflösung zwischen 0,5 und 1 m nur selten zum Einsatz. Eine Ausnahme ist bei Fung und Chan (1994) zu finden. Eine Vielzahl von Beispielen gibt es für die Verwendung rein spektraler Merkmale für die Klassifikation multitemporaler Satellitenbilder mittlerer Auflösung (4 bis 30 m), u.a. in (Gong et al., 2003), (Hall et al., 2003), (Itzerott und Kaden, 2006), (Itzerott und Kaden, 2007), (Marçal und Cunha, 2007) und (Hall et al., 2008).

Fung und Chan (1994) verwenden ein SPOT-Bild für die Klassifikation von 8 *urbanen Klassen*. Als Klassifikationsalgorithmus diente zunächst eine k-means-Klassifikation, mit der 40 spektrale Klassen im Bild gefunden und anschließend zu 8 *urbanen Klassen* zusammengefasst werden. Bei der Evaluation konnte eine *overall accuracy* von 84,4% und ein Kappa-Koeffizient von 81,9% erreicht werden. Dies ist ein sehr gutes Ergebnis, leider decken die Klassen in (Fung und Chan, 1994) nur bedingt die Klassen ab, die Teil der vorliegenden Arbeit sind.

Itzerott und Kaden (2006, 2007) nutzen eine multitemporale Analyse für die Klassifikation landwirtschaftlicher Nutzflächen wie *verschiedene Getreidesorten* und *Feldgras* mittels NDVI-Normkurven. Die NDVI-Normkurven werden für eine objektbasierte Klassifikation mittels Box- bzw. ML-Klassifikationsverfahren verwendet. Datengrundlage sind vier Landsat-Bilder, die innerhalb eines Jahres aufgenommen wurden sowie ein Schlagkataster. Itzerott und Kaden (2006, 2007) konnten mit der spektralen multi-temporalen Analyse eine *overall accuracy* von 65,7% mit dem Box-Klassifikationsverfahren und von 72,8% mit dem ML-Klassifikationsverfahren erreichen. Die Übertragung des Ansatzes innerhalb des Gebietes sowie auf angrenzende Gebiete auf ein anderes Jahr ist nach Angaben der Autoren möglich. In diesem Fall muss vor der Klassifikation eine Einpassung der Aufnahmezeitpunkte in das Normjahr der phänologischen Entwicklung der Kulturen erfolgen, was nicht immer einfach ist, da die Phänologie in unterschiedlichen Gebieten teils erheblich variieren kann. Eine Limitierung des Ansatzes liegt vor, wenn die Anzahl und die Zeitpunkte der Aufnahmen im Anbaujahr nicht ausreichend sind.

Ebenfalls unter Verwendung des NDVI aus einem multitemporalen Datensatz unter Nutzung eines Schlagkatasters bestimmen Marçal und Cunha (2007) *Weingärten*. Bei dem Datensatz handelt es sich um neun SPOT 5 Bilder, aufgenommen in den Jahren 2002, 2003 und 2005 sowie vier Bilder des Satelliten Chris Proba (17 m Bodenauflösung, 18 Kanäle im Bereich rot bis infrarot) aus dem Jahr 2006. Neben den Durchschnitts-NDVI-Werten werden ebenfalls das Minimum, das Maximum, die Standardabweichung und der Median des NDVI-Wertes jedes GIS-Objektes verwendet. Marçal und Cunha (2007) stellen im Abschluss ihres Artikels heraus, dass die abgeleiteten Merkmale aus den NDVI-Bildern geeignet sind, um eine Klassifizierung von Weinanbauflächen durchzuführen. Quantitative Aussagen werden nicht gemacht.

Spektrale Merkmale werden nicht nur für die Klassifikation, sondern auch für die Analyse der Beschaffenheit von Pflanzen genutzt. Exemplarisch sollen hier die Arbeiten von Hall et al. (2003, 2008) vorgestellt werden, die die Lage und weitere Eigenschaften wie Biomasse, Form und Größe von Weinreben in Weinanbauflächen betrachten. Bei der Klassifikation handelt es sich um eine pixelbasierte Klassifikation eines multi-temporalen Bilddatensatzes, bestehend aus drei multispektralen Luftbildern (25 cm Auflösung), aufgenommen zu unterschiedlichen Zeitpunkten innerhalb eines Jahres zu phänologisch wichtigen Zeitpunkten innerhalb der Sommermonate. Der Bewuchs zwischen und unter den Weinreben wurde eine Woche vor der Befliegung entfernt, so dass eine Trennung zwischen Wein und Boden mittels einfachen NDVI-Schwellwerts bestimmt werden kann. Nach der Trennung von Weinrebe und Boden werden die Lage und Ausrichtung der Pflanzreihen bestimmt. Anschließend werden Größe und weitere Eigenschaften der Weinrebe aus dem NDVI abgeleitet. Hall et al. (2003, 2008) standen sehr umfangreiche Vorinformationen wie Abstand zwischen den Pflanzen und zwischen den Pflanzenreihen sowie die Hangneigung und Traubenart zur Verfügung, was für Klassifikationsaufgaben eine Ausnahme darstellt. Nähere Aussagen, wie der Schwellwert zur Bestimmung von Weinrebe und Boden getroffen wurde, werden nicht gemacht. Außerdem ist es nicht immer möglich, die von Hall et al. (2003, 2008) durchgeführten Vorbereitungsschritte vor jeder Befliegung vorzunehmen, speziell das Entfernen von Pflanzen, die keine Weinreben sind.

2.1.3.2. Klassifikation mittels textueller Merkmale

Ausschließlich textuelle Merkmale unter Anwendung von Texturparametern, abgeleitet aus der GLCM, werden in (Smits und Annoni, 1999) verwendet. Ziel dieser Arbeit ist die Aktualisierung eines CLC-GIS-Datensatzes, speziell der Klassen *Siedlung*, *Industrie* und *landwirtschaftliche Fläche*. Zur Klassifikation werden aktuelle Satellitendaten (10 m oder besser), u.a. des IRS-1C- Satelliten (panchromatische Bilder mit einer spektralen Auflösung von 5,8 m) verwendet. Die pixelbasierte Klassifikation mittels MRF wird u.a. mit Hilfe der Haralick-Merkmale Energie, Entropie, Homogenität und Kontrast durchgeführt. Es werden keine Aussagen getroffen, wie die pixelbasierte Analyse auf GIS-Objekte übertragen wird. Die *overall accuracy* für landwirtschaftliche Flächen der pixelbasierten Klassifikation lag bei 98%, wobei eine weitere Unterteilung der landwirtschaftlichen Flächen (wie z.B. in Acker- und Grünland) nicht erfolgte. Dieser Ansatz ist geeignet, landwirtschaftliche Flächen gegenüber nicht landwirtschaftlichen Flächen abzugrenzen, jedoch können keinerlei Aussagen zur Trennung unterschiedlicher landwirtschaftlicher Objektarten gemacht werden.

Ziel von Rengers und Prinz (2009) ist die Klassifikation der CLC-Landnutzungsklassen *Ackerland*, *Wald*, *Wasser*, *Grünland* und *urbane Flächen*. Zunächst wird dafür das Eingabebild in ein Grauwertbild umgewandelt. Nachdem Segmente mittels Schachbrettsegmentierung des Bildes erzeugt wurden, werden Maße wie Komplexität und Texturstärke, abgeleitet aus der NGTDM, für jedes Segment berechnet. Das Ergebnis zeigt, dass die Trennung von Acker- und Grünland nicht möglich ist, die vereinte Klasse Acker- und Grünland dafür aber sicher von allen anderen Klassen getrennt werden kann. Warum das Eingangsbild in ein Grauwertbild umgewandelt wird und damit wertvolle Farbinformationen bei der Klassifikation nicht berücksichtigt werden, wird nicht geklärt. Es wird ebenfalls nicht erläutert, wie das Grauwertbild erstellt wird. Außerdem werden keine Angaben über die Größe der Segmente gemacht.

2.1.3.3. Klassifikation mittels struktureller Merkmale

Trias-Sanz (2006) benutzt das Semi-Variogramm für die objektbasierte Klassifikation von GIS-Objekten eines Schlagkatasters mit den Objektklassen *Wald*, *bepflanzte und unbepflanzte Ackerlandflächen*, *Plantagen* sowie *Weingärten* unter Nutzung eines RGB Luftbildes mit einer Bodenauflösung von 50 cm. Diese Objektklassen unterscheiden sich durch ihre Hauptbewirtschaftungsrichtungen und die unterschiedlichen Abstände der Bewirtschaftungsrichtungen (Periodizität). So besitzt *Wald* keine Hauptbewirtschaftungsrichtung, *Ackerlandflächen* eine und *Plantagen* und *Weinbauf Flächen* zwei, wobei sich die beiden Klassen durch ihre Periodizität unterscheiden. Die Erstellung des Semi-Variogramms erfolgt für den inneren Bereich jedes GIS-Objektes. Nachdem im Semi-Variogramm Extremwerte berechnet wurden, wird mittels der Extrema untersucht, ob keine oder mindestens eine Hauptrichtung vorliegt. Keine Hauptrichtung liegt vor, wenn im Semi-Variogramm höhere Werte in der Nähe des Ursprungs auftreten, verglichen zu denen, die weiter vom Ursprung entfernt liegen. Liegt mindestens eine Hauptrichtung vor, wird das Semi-Variogramm interpoliert und in einen geeigneten Akkumulationsraum (r , α) überführt, der nicht näher bezeichnet wird. r und α dieses Akkumulationsraums sind dabei vergleichbar mit den Größen d und φ des Hough-Raums. Warum das Semi-Variogramm interpoliert wird, ist nicht bekannt. Gibt es parallele Strukturen mit einer einheitlichen Hauptrichtung im Bild, so entstehen mehrere Maxima mit gleichem α im Akkumulationsraum. Die Häufigkeit, mit der eine Richtung vorkommt, kann mit Hilfe eines Histogramms, abgeleitet aus dem Akkumulationsraum, bestimmt werden. Das Maximum dieses Histogramms entspricht der Hauptrichtung im Bild. Die Periodizität wird ermittelt, indem entlang der bestimmten Hauptrichtung im Semi-Variogramm ein Schnitt durchgeführt und ein weiteres Histogramm abgeleitet wird. In dem so entstandenen zweiten Histogramm werden zunächst die Extrema mittels Wasserscheidentransformation bestimmt. Aus den so bestimmten Extrema wird ein Stärkemaß abgeleitet, mit dessen Hilfe bestimmt wird, ob im Histogramm mehrere signifikante Maxima enthalten sind. Ist dies der Fall, kann aus diesen signifikanten Maxima die Periodizität zwischen den im Bild vorhandenen parallelen Strukturen bestimmt werden. Das Vorhandensein einer Hauptrichtung und eine eventuell vorhandene Periodizität werden für eine regelbasierte Klassifikation benutzt, bei der eine *overall accuracy* von 95,5% erreicht wird. Der in (Trias-Sanz, 2006) beschriebene Ansatz kann für die Trennung einer großen Anzahl landwirtschaftlicher Nutzflächen verwendet werden. Schwierigkeiten können auftreten, wenn es sich um eine nicht homogene Fläche im GIS-Objekt handelt, z.B. sich der Abstand der parallelen Strukturen oder die Hauptbewirtschaftungsrichtung selber ändert. Außerdem ist eine weitere Voraussetzung für diesen Algorithmus das Vorhandensein eines Schlagkatasters.

Warner und Steinmaus (2005) verwenden die Autokorrelation für eine pixelbasierte Klassifikation von *Weingärten* und *Plantagen* in IKONOS-Bildern. Nachdem das Eingabebild geglättet und zu kleine GIS-Objekte entfernt wurden, wird eine quadratische, für die Objektklasse charakteristische Mustermatrix (ein Bildausschnitt, hier Kernel genannt) innerhalb eines GIS-Objektes gewählt. Die Größe des Kernels wird vom Benutzer vorgegeben; Aussagen, wie der Kernel gewählt wird, werden nicht gemacht. Der Kernel wird in der lokalen Nachbarschaft (Suchmatrix) entlang der vier Hauptrichtungen sowie beider Diagonalen bewegt und pro Pixel und Richtung wird der Autokorrelationswert bestimmt. Wenn die Entfernungen der Korrelationsmaxima ungefähr gleichabständig sind, liegt ein Bildausschnitt mit gleichmäßigen Strukturen vor, was ein Hinweis auf *Weingärten* oder *Plantagen* ist. Ein Weingarten- oder Plantagenpixel liegt vor, wenn mindestens 5 Maxima vorhanden sind und eine Varianz zwischen den Abständen der Maxima nicht größer als ein zuvor definierter Schwellenwert vorliegt. Evaluiert wurde der Ansatz in einem Gebiet, in dem neben *Plantagen* auch *Weingärten*, *Getreide* und *Grünland* zu finden waren. Als Kernelgröße wurden 21 x 21 Pixel und als lokale Nachbarschaft 25 x 25 Pixel festgelegt. Es konnte eine *overall accuracy* von 95,4% und ein Kappa-Koeffizient von 90% erreicht werden, wobei die mittlere *user's accuracy* über alle Objektklassen bei 95,3% und die mittlere *producer's accuracy* über alle Objektklassen bei 91,2% lag. Das Ergebnis wird mit

einem Klassifikationsergebnis verglichen, das mittels einer ML-Klassifikation von Haralick-Merkmalen entstand. Dabei wurden ca. 25% der Objekte als Trainingsdaten benutzt. Die so erzielte *overall accuracy* lag bei 86,5%, der Kappa-Koeffizient bei 70,1%, die mittlere *user's accuracy* über alle Objektklassen bei 74,4% und die mittlere *producer's accuracy* über alle Objektklassen bei 75,9%. Das Verfahren mittels Autokorrelation ist besonders bei großen Gebieten mit regelmäßig wiederkehrenden Strukturen geeignet. Als positiv wird in der Veröffentlichung herausgestellt, dass neben der Klassifikation auch zusätzliche Merkmale wie die Pflanzrichtung und der Pflanzabstand ermittelt werden können, wobei dies mittels anderer Techniken wie die des Semi-Variogramms auch möglich ist. Ein weiterer Vorteil ist, dass ein Training nicht erforderlich ist. Ein Nachteil dabei ist, dass die Gleichabständigkeit der Bäume/Pflanzen für diesen Algorithmus eine Voraussetzung ist und viele Parameter vom Benutzer eingestellt werden müssen (z.B. die Größe des Kernels).

Wassenaar et al. (2002) wenden eine Fouriertransformation zum Suchen von räumlichen Strukturen in GIS-Objekten eines Schlagkatasters zur objektbasierten Klassifikation von *Weingärten*, *Plantagen* sowie *Getreide/Brachflächen* an. Zur Verfügung stehen Luftbilder mit 25 cm Bodenauflösung, wobei für die Analysen auf Grund des guten Kontrastes zwischen Vegetation und Nicht-Vegetation nur der Rotkanal verwendet wird. Die Fouriertransformation wird nur auf den inneren Bereich des GIS-Objektes angewendet. Anschließend werden im entstandenen Amplitudenbild Maxima gesucht, die Suche wird unter Anwendung von Vorwissen über die Abstände von Pflanzreihen eingeschränkt. Mittels der detektierten Maxima ist die Bestimmung der Hauptbewirtschaftungsrichtung sowie der Abstände zwischen den Pflanzenreihen möglich. Beide Merkmale werden für die Klassifikation mittels eines baumförmigen Klassifikationsverfahrens genutzt. Im Testgebiet, das zu 70% aus *Weingärten* und u.a. auch aus *Getreidefeldern* (Winterweizen und Raps) und kleinen *Plantagen* besteht, konnten alle *Weingärten* erfolgreich detektiert werden. Voraussetzung für den in (Wassenaar et al., 2002) vorgestellten Algorithmus ist, dass die Abstände der Bewirtschaftungsspuren im GIS-Objekt konstant und gegeben sind, was im Fall der in dieser Veröffentlichung zu detektierenden Weinanbauflächen gilt. Bei Ackerland hingegen können die Abstände der Pflanzreihen bzw. Bewirtschaftungsspuren je nach Fruchtart und verwendeter Bewirtschaftungsmaschinen stark schwanken und stehen im Normalfall nicht zur Verfügung.

Der Schwerpunkt der Arbeit von Delenne et al. (2007) ist ebenfalls die Klassifikation von *Weingärten* und *Nicht-Weingärten* sowie die Bestimmung der Bewirtschaftungsrichtung und des Abstandes zwischen den Pflanzreihen. Zur Verfügung stand ein Luftbild mit 50 cm Bodenauflösung. Neben der Anwendung der Fouriertransformation wird vergleichend auch eine pixelbasierte Klassifikation mit Hilfe des Haralick-Merkmals Kontrast getestet. Die Anwendung der Fouriertransformation ist angelehnt an den Algorithmus von Wassenaar et al. (2002). Das Haralick-Merkmal Kontrast wurde nach der Analyse einer Testreihe verschiedener Haralick-Merkmale als am geeignetsten ermittelt und ausgewählt. Für den Algorithmus wurde genauso wie in (Wassenaar et al., 2002) nur der Rotkanal verwendet. Für die Berechnung der Merkmale wurde auch die lokale Nachbarschaft (31 x 31 Pixel) um einen Pixel betrachtet. Ergebnis der Analysen ist ein „Weinindex“, der die Bewirtschaftungsrichtung sowie die Periodizität im Bildausschnitt berücksichtigt. Da das Ziel die Bewertung von GIS-Objekten ist, ist der finale Schritt, das pixelbasierte Klassifikationsergebnis auf ein GIS-Objekt zu übertragen. Ein Objekt muss mindestens zu 75% aus der Klasse *Weingarten* bestehen, um als Weingarten bzw. zu mindest 75% aus der Klasse *Nicht-Weingarten*, um als Nicht-Weingarten klassifiziert zu werden. Falls keiner der beiden Fälle eintritt, wird das Objekt als nicht klassifizierbar markiert. Der Algorithmus wurde in einem Testgebiet evaluiert, das zu 70% aus Weingärten bestand. Bei einer Klassifikation mittels Fouriertransformation konnten die Ergebnisse gegenüber einer Klassifikation mittels des Haralick-Merkmals Kontrast verbessert werden. Die *overall accuracy* mittels Haralick-Merkmal beträgt 70,8%, die mittels Fouriertransformation 86,6%. Beide Algorithmen erbringen bessere Ergebnisse für Objekte der Klasse *Weingarten* als für Objekte der Klasse *Nicht-Weingarten*. Für den Ansatz nach Delenne et al. (2007) gelten ähnliche Kritikpunkte wie für den Ansatz nach Wassenaar et al. (2002), weil er im Wesentlichen auf letzteren basiert.

Ein anderes Verfahren zur Bestimmung von *Weinanbauflächen* und speziell für die Bestimmung der Bewirtschaftungsrichtung und des Abstandes der Pflanzreihen wird in (Chanussot et al., 2005) vorgestellt. Auch hier stehen RGB Luftbilder (mit unbekannter Auflösung) zur Verfügung, wobei das Differenzbild aus rotem und grünem Kanal erzeugt als Eingangsbild für den Algorithmus dient. Die Fouriertransformation wird auf das Eingabebild eines jeden GIS-Objektes eines Schlagkatasters angewendet. Im jeweils entstandenen Amplitudenbild sind Maxima orthogonal zur Bewirtschaftungsrichtung der Weinpflanzen zu

erkennen, die sich wie Streifen durch das Spektrum ziehen. Um diese Streifen zuverlässig einer Bewirtschaftungsrichtung zuordnen zu können, wird auf das Amplitudenbild eine spezielle Houghtransformation angewendet, wobei keine weiteren Angaben über Vorverarbeitung des Amplitudenbildes gemacht werden. Im Hough-Raum sind Maxima zu erkennen, die mit den Bewirtschaftungsrichtungen korrespondieren. Neben der Hauptbewirtschaftungsrichtung wird auch der Abstand zwischen den Reihen der Weinreben aus dem Hough-Raum abgeleitet. Die beschriebenen Analyseschritte sind Grundlage für die Detektion toter Weinpflanzen. Detaillierte Angaben über den Erfolg der Anwendung dieses Algorithmus zur Klassifikation bzw. der Detektion toter Weinpflanzen sind in der Veröffentlichung nicht gegeben, genauso wenig wie Angaben über verwendete Parameter oder eine Analyse der Zuverlässigkeit.

Eine spezielle Fouriertransformation findet Anwendung in den Arbeiten von Rabatel et al. (2008) und Delenne et al. (2008), in denen *Weingärten* klassifiziert sowie deren Grenzen detektiert werden. Es stehen RGB Luftbilder mit einer Auflösung von 50 cm zur Verfügung, von denen nur der Rotkanal verwendet wird. Das Eingabebild wird mittels Schachbrettsegmentierung in 500 x 500 Pixel große Segmente zerlegt. Die so entstandenen Bildsegmente können verschiedene *Weingärten* enthalten, die es zu klassifizieren und zu detektieren gilt. Die Analyse beschränkt sich auf den inneren Bereich jedes Segmentes, um Einflüsse des Randes bei der Fouriertransformation zu vermeiden. Die Hauptidee besteht in der Isolierung der einzelnen *Weingärten* durch Selektion der korrespondierenden Frequenz eines *Weingartens* im Amplitudenbild. Detaillierte Aussagen über den Erfolg der Klassifikation werden nicht gemacht. Was geschieht, wenn ein *Weingarten* größer als das Segment ist, wird ebenfalls nicht geklärt.

2.1.3.4. Klassifikation mittels spektraler und textueller Merkmale

Von Haralick (1973) werden spektrale Merkmale (Mittelwert und Standardabweichung der Grauwerte der Kanäle) sowie textuelle Merkmale (Energie, Kontrast, Korrelation und Entropie, abgeleitet aus der GLCM), zur Klassifikation von verschiedenen Landbedeckungsklassen (*Küstenwald, forstliche Nutzfläche, Grünland, urbane Gebiete, kleine und große bewässerte Felder und Wasser*) genutzt. Dafür stand ein vierkanaliges multispektrales Satellitenbild von geringer geometrischer Auflösung (ca. 80 m Bodenauflösung) zur Verfügung. Die Klassifizierung wurde pixelbasiert unter Berücksichtigung der lokalen Umgebung der Pixel durchgeführt. Als Klassifikator wurde das *LDF-Verfahren* verwendet. Bei der Klassifikation konnte eine *overall accuracy* von 83,5% erreicht werden, wobei die *user's accuracy* von Grünland bei 77,3% lag. Dies stellt eine Steigerung der *overall accuracy* gegenüber der Klassifikation mit nur spektralen Merkmalen (74-77%) dar. Da die Textur skalenabhängig ist, ist unklar, ob die in (Haralick et al., 1973) erreichten Genauigkeiten auch für Bilder mit einer geometrischen Genauigkeit von nur 0,5 bis 1 m erreicht werden können, wie sie in dieser Arbeit vorliegen.

Gong et al. (2003) verwenden die Haralick-Merkmale Mittelwert, Varianz, Homogenität, Kontrast, Entropie, Korrelation, Unähnlichkeit und Trägheitsmoment, abgeleitet aus den NIR-Kanal sowie die spektralen Merkmale NDVI, „greenness“ und „brightness“, abgeleitet aus den Kanälen eines vierkanaligen multispektralen Luftbildes mit einer Bodenauflösung von 2 m für die Unterscheidung der Klassen *Weingärten* und *Nicht-Weingärten* mittels einer pixelbasierten ML-Klassifikation. Unter Verwendung der Kombination verschiedener Merkmale konnte eine *overall accuracy* von maximal 81% erreicht werden, wobei diese durch verschiedene Nachbearbeitungsschritte bis auf durchschnittlich 87% verbessert werden konnte. Schwierigkeiten der zuverlässigen Klassifikation von *Weingärten* konnten bei diesem Ansatz bei sehr niedrigen Weinpflanzen festgestellt werden, die häufig mit Gras verwechselt werden. Die textuellen Merkmale, die in dieser Veröffentlichung gewählt werden, reichen daher alleine nicht aus, um *Weingärten* von *Grünland* zuverlässig trennen zu können. Es ist zu bezweifeln, dass mit dieser Methode eine zuverlässige Trennung von Acker- und Grünland möglich ist.

Das Hauptziel von Hoberg und Rottensteiner (2010) ist der Nachweis, dass durch die Berücksichtigung von Kontextwissen das Klassifikationsergebnis verbessert werden kann. Als Beispiel diente die Klassifikation von Siedlung mittels eines RGB-IKONOS-Bildes. Als Merkmale werden spektrale und textuelle Merkmale verwendet, die innerhalb zweier unterschiedlich großer lokaler Nachbarschaft eines Pixels berechnet werden, die als Skalen bezeichnet werden. Als textuelle Merkmale werden gradientenbasierte Merkmale benutzt. Um diese zu bestimmen, werden zunächst der Gradientenwert und die Richtung für jeden Pixel innerhalb eines der Skalen entsprechenden Fenster berechnet. Basierend darauf wird für jede Skala ein gewichtetes Histogramm der Gradientenrichtungen erstellt. Aus diesem Histogramm werden Texturmaße abgeleitet, dazu gehören der Mittelwert und die Varianz des Histogramms sowie die Anzahl der Maxima im

Histogramm, die über dem Mittelwert liegen. Diese Merkmale liegen für jede Skala vor. Zu den spektralen Merkmalen gehört die Varianz des Farbtones nach Transformation des RGB-Bildes in einen HSV-Farbraum. Dies erfolgt wiederum für jede Skala. Die Klassifikation erfolgte mittels des CRF-Klassifikationsverfahrens, welches eine *producer's accuracy* von bis zu 94,4% und eine *user's accuracy* von bis zu 91,6% für die Siedlungsklasse erreichen kann, während mit dem ML-Verfahren für diese Klasse eine *producer's accuracy* von bis zu 90,7% und eine *user's accuracy* von bis zu 75,8% für die Siedlungsklasse erreicht wird. Dieser Arbeit zeigt, dass die Einbeziehung von Kontextwissen z.B. durch die Wahl des Klassifikationsverfahrens den Erfolg einer Klassifikation steigern kann. Aus diesem Grund lag der Fokus nicht auf der Eignung der Merkmale für die Klassifikation, im Gegenteil, sie wurden einfach gehalten, sondern in der Eignung des Klassifikationsverfahrens. Die Merkmale wurden gewählt, um die Klasse *Siedlung* klassifizieren zu können. Der Erfolg dieser Merkmale für die Klassifikation von landwirtschaftlichen Nutzflächen ist offen und kann nicht abschließend bewertet werden.

2.1.3.5. Klassifikation mittels spektraler und struktureller Merkmale

In (Ranchin et al., 2001) erfolgt eine pixelbasierte Klassifikation von Flächen in die Klassen *Weingärten* und *Nicht-Weingärten* mittels Falschfarbeninfrarotbild (IRRG) mit hoher Auflösung (50 cm Bodenauflösung). Für die Klassifikationsaufgabe werden sowohl spektrale Merkmale, abgeleitet aus dem grünen Kanal, als auch strukturelle Merkmale verwendet. Die strukturellen Merkmale basieren auf der Detektion von Linien mittels Fouriertransformation. Die Liniendetektion beruht auf der Interpretation eines Histogramms, das aus dem Amplitudenbild abgeleitet wird. Da das Ziel die Bewertung von GIS-Objekten ist, ist der finale Schritt, das pixelbasierte Klassifikationsergebnis auf ein GIS-Objekt zu übertragen, wobei genauso wie bei Delenne et al. (2007) vorgegangen wird. Mit dem Algorithmus von Ranchin et al. (2001) konnte eine *overall accuracy* von 82% erreicht werden, wobei *Weingärten* etwas zuverlässiger als *Nicht-Weingärten* detektiert werden konnten (*user's accuracy* der Klasse *Weingarten*: 96%; der Klasse *Nicht-Weingarten* 68%; *producer's accuracy* der Klasse *Weingarten*: 74 %, der Klasse *Nicht-Weingarten* 95%). Weitere Analysen ergaben, dass bei der Verwendung von Bilddaten mit einer Bodenauflösung größer als 50 cm der Erfolg der Klassifikationsergebnisse sank. Bei der Übertragung des Ansatzes auf andere Weinanbauflächen in Italien (1 m Bodenauflösung), Spanien (50 cm Bodenauflösung) und Frankreich (60 cm Bodenauflösung) wurden weniger gute Ergebnisse erreicht. Zur Anwendung des Algorithmus auf diese Gebiete war neben einer Anpassung des Algorithmus auf die neuen Bilddaten auch eine Anpassung auf abweichende Eigenschaften der Pflanzenanbautechnik (u.a. Abstand der Pflanzreihen) notwendig.

2.1.3.6. Klassifikation mittels spektraler und geometrischer Merkmale

Peled und Gilichinsky (2010) verwenden neben den Merkmalen, abgeleitet aus dem NDVI (Mittelwert und Standardabweichung), auch geometrische Merkmale (Fläche, Kompaktheit – jeweils mit Mittelwert und Standardabweichung) zur Klassifizierung von verschiedenen *landwirtschaftlichen Flächen*, *Wald* und *Wasser*. Zunächst findet eine Segmentierung innerhalb der GIS-Objekte statt. Die entstandenen Segmente werden anschließend an Hand geometrischer und spektraler Merkmale klassifiziert. Wenn ein GIS-Objekt aus mindestens 80% der im GIS enthaltenen Objektklasse besteht, wird das GIS-Objekt als korrekt angenommen. Insgesamt wurde eine *overall accuracy* von 83,6% und ein Kappa-Koeffizient von 71% erreicht, wobei nur 62,0% der landwirtschaftlichen Nutzflächen richtig klassifiziert werden konnten. Dieses Ergebnis ist daher weniger zufriedenstellend.

2.1.3.7. Klassifikationsverfahren mit Merkmalen aus mehr als zwei Merkmalsgruppen

In den Arbeiten von Ruiz et al. (2004, 2007), Balaguer et al. (2010) und Recio et al. (2011) werden Merkmale aus allen Merkmalsgruppen verwendet. Alle führen eine objektbasierte Klassifikation von unterschiedlichen landwirtschaftlichen Objektklassen wie *Ackerland*, *Strauchgewächs*, *Wald* und *Zitrusplantage* mittels multispektraler Luftbilder mit einer Auflösung von 50 cm durch. Bei dem verwendeten GIS handelt es sich immer um ein Schlagkataster. Bei allen Veröffentlichungen werden Merkmale aus allen Merkmalsgruppen verwendet, die in Tabelle 2 zusammengefasst sind. Eine detaillierte Beschreibung der textuellen Merkmale ist (Ruiz et al., 2004), der Merkmale abgeleitet aus dem Semi-Variogramm (Balaguer et al., 2010), weitere strukturelle Merkmale (Ruiz et al., 2007) und (Balaguer et al., 2010) sowie weiterer Merkmale (Hermosilla et al., 2010) und (Recio et al., 2011) zu entnehmen. Die Klassifikation der Objekte an Hand der Merkmale wird mit Hilfe des C5-Klassifikationsverfahrens durchgeführt.

Merkmalsgruppe	Verfahren/abgeleitete Merkmale
Spektrale	Je GIS-Objekt pro Farbband plus NDVI: • Mittelwert • Standardabweichung • Minimum • Median • Maximum • Bandweite (Maximum – Minimum) • Summe
Texturelle	• Haralickmerkmale (Kontrast, Einheitlichkeit, Entropie, Varianz, Kovarianz, Inverse Difference Moment und Korrelation) • Farbhistogramm (Schiefheit und Kurtosis) • Kantenfaktor (edgeness factor) nach (Sutton und Hall, 1972) • Dichte der Kanten in der Nachbarschaft
Strukturelle	• Semi-Variogramm, s. (Balaguer et al., 2010) • Houghtransformation, s. (Ruiz et al., 2007) • Fouriertransformation, s. (Ruiz et al., 2007)
Weitere	• Kompaktheit des GIS-Objektes nach (Bogaert et al., 2000) • Formindex und fraktale Dimension nach (Krummel et al., 1987) • Fläche und Umfang des GIS-Objektes • Merkmale aus weiteren Quellen, u.a. (Recio et al., 2011)

Tabelle 2: Merkmale benutzt in (Ruiz et al., 2004), (Ruiz et al., 2007), (Balaguer et al., 2010) und (Recio et al., 2011).

Ergebnisse sind u.a. (Balaguer et al., 2010) zu entnehmen, wo spektrale Merkmale, strukturelle Merkmale, abgeleitet aus einem Semi-Variogramm, und texturelle Merkmale (Haralick-Merkmale) verwendet werden. In (Balaguer et al., 2010) wird gezeigt, dass die *producer's accuracy* unter Verwendung des Semi-Variogramms und spektraler Merkmale höher ist als die, die mittels einer Klassifikation mit Haralick-Merkmalen und spektraler Merkmale erzielt wird. Es konnte mittels des Semi-Variogramms eine mittlere *producer's accuracy* von 90,7% ohne Verwendung zusätzlicher spektraler Merkmale und eine mittlere *producer's accuracy* von 96,0% mit zusätzlicher spektraler Merkmale (NDVI) erreicht werden. Es kann zu Problemen kommen, wenn nicht homogene Flächen vorliegen, also sich z.B. Abstand oder Hauptbewirtschaftungsrichtung selbst ändern.

Ruiz et al. (2007) konnten unter Verwendung von spektralen, texturellen und strukturellen Merkmalen eine *overall accuracy* von 89,9% erreichen. Zu Fehlklassifikationen kam es, wenn das Objekt nicht dem Modell entsprach (z.B. Olivenbäume, die nicht in einem regelmäßigen Muster angeordnet sind) oder wenn die Bildauflösung zu gering war (z.B. wenn junge Orangenbäume nicht extrahiert werden konnten). Zu Fehlern bei der Klassifikation kommt es auch, wenn kleine Teile der Objekte mit einer anderen Klasse bedeckt sind, die nicht detektiert werden können, z.B. eine Siedlung, die in eine Plantage reicht.

Die Arbeiten von Ruiz et al. (2004, 2007), Balaguer et al. (2010) und Recio et al. (2011) zeigen den Vorteil der Kombination von Merkmalen unterschiedlicher Merkmalsgruppen. Schwerpunkt dieser Arbeiten ist eine Mehrklassenklassifikation. Die verwendeten Merkmale sind daher darauf ausgelegt, mehrere Klassen zuverlässig voneinander zu trennen. Eine Untersuchung, ob diese Merkmale auch auf das Zwei-Klassenproblem Acker- und Grünland für Bilder mit einer geometrischen Auflösung von 0,5 bis 1 m zutreffen, ist bisher nicht bekannt. Hinzu kommt, dass der Ansatz von Ruiz et al. (2004, 2007), Balaguer et al. (2010) und Recio et al. (2011) ein Schlagkataster benötigt. Außerdem können z.B. Gebäude in einem GIS-Objekt mit landwirtschaftlicher Nutzung nicht detektiert werden, obwohl dies wünschenswert ist.

2.2. Verifikation von GIS-Objekten

Wie in Kapitel 1 definiert, ist die Verifikation das Finden aller Abweichungen zwischen GIS und Bilddaten, also das Finden von Objekten im GIS, die in der Realität nicht existieren, sich geometrisch geändert haben bzw. in den GIS-Daten falsch repräsentiert sind. Arbeiten mit dem Schwerpunkt der Verifikation sind (Walter, 2004), (Busch et al., 2004), (Busch et al., 2005), (Helmholz et al., 2007b) und (Helmholz et al., 2010b). Des Weiteren wurden in Kapitel 1 die *change detection* und das *update* als Verfahren zum Qualitätsmanagement von GIS-Daten vorgestellt. Arbeiten mit dem Schwerpunkt der *change detection* sind u.a. (Knudsen und Olsen, 2003), (Olsen, 2004), (Carleer und Wolff, 2008) und (Leignel et al., 2010); Arbeiten mit dem Fokus auf das *update* sind u.a. (Gerke et al., 2003), (Grote und Heipke, 2008) und (Rottensteiner, 2008).

Die meisten der Veröffentlichungen, die sich mit Verifikation, *change detection* oder *update* beschäftigen, befassen sich nicht ausschließlich mit landwirtschaftlichen Nutzflächen. Daher soll zunächst auf die Überprüfung von GIS allgemein eingegangen werden. Bei der Betrachtung der Literatur lässt sich feststellen, dass es sich um ein Thema von internationalem Interesse handelt. So gibt es Beiträge aus Australien (Rottensteiner, 2008), Belgien (Carleer und Wolff, 2008), (Leignel et al., 2010), Dänemark (Knudsen und Olsen, 2003), Deutschland (Straub et al., 2000), (Busch et al., 2004), (Walter, 2004), (Busch et al., 2005), (Helmholz et al., 2010b), (Krickel, 2010), Frankreich (Champion et al., 2010), Großbritannien (Holland et al., 2006), Italien (Gianinetto, 2008) und der Schweiz (Eidenbenz et al., 2000). In allen genannten Publikationen besteht bereits ein GIS – das Erzeugen eines GIS ist nicht notwendig, vielmehr ist die Verifikation bzw. das *update* im GIS existierender Daten von Relevanz. Alle Ansätze haben den Vergleich des GIS mit aktuellen Bilddaten gemeinsam, wobei die dazu verwendeten Bilddaten stark variieren. Die am häufigsten genutzten Bilddaten sind multispektrale Luft- und Satellitendaten (Straub et al., 2000), (Eidenbenz et al., 2000), (Knudsen und Olsen, 2003), (Busch et al., 2004), (Olsen, 2004), (Walter, 2004), (Busch et al., 2005), (Holland et al., 2006), (Helmholz et al., 2010b), (Carleer und Wolff, 2008), (Gianinetto, 2008), (Poulain et al., 2008), (Champion et al., 2010), (Hanson und Wolff, 2010), (Krickel, 2010), (Leignel et al., 2010), (Poulain et al., 2011), aber auch digitale Oberflächenmodelle (Eidenbenz et al., 2000), (Rottensteiner, 2008), (Akca et al., 2010), (Champion et al., 2010), (Matikainen et al., 2010), SAR-Daten (Poulain et al., 2008), (Hanson und Wolff, 2010), (Poulain et al., 2011) oder weitere Informationen wie z.B. Bauunterlagen (Krickel, 2010). Ein weiterer Unterschied zwischen den Verifikationsansätzen ist die Nutzung von Vorinformationen, die durch das GIS gegeben sind. Da bei dieser Arbeit ausschließlich monotemporale multispektrale Bilder zu Verfügung stehen, wird in diesem Abschnitt nur auf Verifikationsverfahren eingegangen, die sich mit der automatischen oder semi-automatischen Verifikation eines existierenden GIS unter der Nutzung von ausschließlich monotemporalen multispektralen Bildern beschäftigt. Ein Schwerpunkt, der verstärkt betrachtet wird, ist die Einbindung des GIS in den Verifikationsprozess.

Walter (2004) benutzt das Wissen, das durch das GIS gegeben ist, auf drei Arten. Zunächst werden die Grenzen der GIS-Objekte für eine objektbasierte Klassifikation mittels ML-Klassifikationsverfahren benutzt. Dabei werden die Trainingsdaten für die überwachte Klassifikation automatisch aus dem GIS abgeleitet, da die Annahme besteht, dass die GIS-Daten zum größten Teil korrekt sind. Für den Klassifikationsprozess werden nur spektrale Merkmale, abgeleitet aus den Bildkanälen eines multispektralen IKONOS-Bildes, verwendet. Die Objektklassen, die dabei unterschieden werden, sind *Wasser*, *Siedlung*, *Wald* und *Grünland*. Anschließend wird das GIS zum Vergleich des Klassifikationsergebnisses mit den GIS-Daten benutzt, um Fehler im GIS-Datensatz zu identifizieren. Die Grenzen des GIS werden nicht bewertet oder aktualisiert. Bei dem Ansatz von Walter (2004) handelt es sich um einen Verifikationsansatz.

Dieselbe Strategie der Nutzung des GIS während der Klassifikation und des Vergleiches der Klassifikationsergebnisse mit dem GIS wird auch in Carleer und Wolff (2008) verwendet, wobei die Grenzen des GIS-Objektes aus der GIS-Datenbank nur als initiale Grenzen verwendet werden. Bevor die Klassifikation durchgeführt wird, wird ein Segmentierungsverfahren auf den inneren Bereich eines GIS-Objektes angewendet. Die entstandenen Segmente werden klassifiziert und die Klassifikationsergebnisse abschließend mit dem GIS verglichen. Ziel der Publikation ist die *change detection* der Klassen *Gebäude*, *Straßen*, *Flächen mit und ohne Vegetation* sowie *Wasser*. Unter Verwendung eines Quick-Bird-Bildes (panchromatisch und multispektrale Kanäle) wird eine *overall change detection accuracy* von 92,2% erreicht.

Ein ähnlicher Ansatz wird in (Leignel et al., 2010) verfolgt. Nachdem eine Segmentierung des inneren Bereiches eines GIS-Objektes mittels verschiedener Techniken wie dem Wasserscheidenverfahren, Graph Cuts oder dem Mean-Shift-Verfahren durchgeführt wurde, werden die Segmente mittels textueller und spektraler Merkmale klassifiziert und anschließend wird das Klassifikationsergebnis mit dem GIS verglichen. Ziel des Verfahrens ist die *change detection* von Veränderungen bei künstlichen Objekten. Tests werden mit IKONOS-Bildern (panchromatisch und multispektrale Kanäle) durchgeführt, eine Evaluation des Ansatzes von Leignel et al. (2010) wird aber leider nicht gegeben.

Ähnlich wie in (Walter, 2004) handelt es sich bei den Arbeiten von Busch et al. (2004, 2005) und Helmholz et al. (2007b, 2010b) um einen reinen Verifikationsansatz. Dieses Verfahren ist im Gegensatz zu allen andern Arbeiten modular aufgebaut und in einem wissensbasierten Bildanalysesystem integriert, womit die Möglichkeit des flexiblen Einbindens verschiedener spezialisierter Verfahren für die Klassifikation und

anschließende Verifikation einer großen Anzahl verschiedener Klassen (u.a. *Siedlung, Industrie, Wald, Ackerland/Grünland, Plantagen, Straßen, Flüsse* und *Schienenwege*); basierend auf unterschiedlichen Bilddaten (Luftbilder, niedrig und hochaufgelöste Satellitenbilder), gegeben ist. Die Klassifikation kann abhängig vom gewählten Ansatz spezialisiert auf die Objektklasse pixel- oder objektbasiert durchgeführt werden. Das verwendete GIS wird bei der Klassifikation zum Teil zum automatischen Trainieren überwachter Klassifikationsverfahren verwendet und zum Teil werden die GIS-Grenzen für objektbasierte Klassifikationsverfahren genutzt. Mit dem Verfahren, vorgestellt in (Busch et al., 2004), (Busch et al., 2005), (Helmholz et al. 2007b) und (Helmholz et al., 2010b), konnte für unterschiedliche Verifikationsaufgaben immer eine *overall accuracy* von über 70% erreicht werden.

Ein komplett anderer Ansatz für die Einbindung des GIS in den Verifikationsprozess ist in (Knudsen und Olsen, 2003) sowie in (Olsen, 2004) zu finden. Beide fokussieren auf die *change detection* bei *Gebäuden*. In einem ersten Schritt werden multispektrale Bilddaten mit dem vektorbasierten GIS fusioniert, indem das GIS in ein Rasterbild überführt wird. Anschließend wird die Gebäudeklasse des rasterisierten GIS mittels k-means Klassifikation in mehrere spektral homogene Gruppen geteilt. Diese Gruppen werden für das Training eines überwachten Klassifikationsverfahrens zur Detektion von Änderungen genutzt. Die so ermittelten Änderungen werden anschließend für die Aktualisierung des GIS verwendet. Die Evaluation zeigt, dass alle Änderungen erfolgreich ermittelt werden konnten, aber die Anzahl der falsch detektierten Änderungen gleichzeitig sehr hoch ist.

Dieser kurze Literaturüberblick zeigt, dass es eine Vielzahl von Ansätzen für Verifikation, *change detection* und *update* von GIS-Daten gibt. Der Kern jedes Verifikationsprozesses ist eine erfolgreiche Klassifikation. *Gebäude, Straßen* und *Vegetation (Bäume)* sind oft Gegenstand von *change detection* und *update*- Verfahren, während bei flächenhaften Geodaten (Landbedeckungsklassen wie *Wasser, Siedlung, Industrie, Wald, Ackerland/Grünland* und *Plantagen*) die Verifikation im Vordergrund steht. Es gibt unterschiedliche Wege GIS-Daten in den Prozess des *updates*, der *change detection* oder der Verifikation einzubinden, wobei die Verwendung des GIS für eine objektbasierte Klassifikation, für das Trainieren überwachter Klassifikationsansätze sowie für den Vergleich der Klassifikationsergebnisse mit dem GIS am häufigsten zu finden sind. Obwohl die gemeinsame Klasse *Acker- und Grünland* bereits Inhalt von Verifikationsverfahren war, wurde eine getrennte Verifikation beider Objektarten bisher nicht durchgeführt.

2.3. Segmentierung von Schlägen in GIS-Objekten

In der Vergangenheit wurden viele Segmentierungsverfahren vorgestellt, z.B. (Gonzalez und Woods, 2002), (Förstner, 1994). Die Anwendbarkeit eines Segmentierungsverfahrens hängt von vielen Faktoren ab. Ein wesentlicher Faktor ist die Feinheit der Segmentierung, die erreicht werden soll. Überblicke über verschiedene Segmentierungsalgorithmen sind in (Eckstein, 1996) und in (Lu und Wenig, 2007) gegeben.

Bei der Segmentierung von Schlägen in GIS-Objekten der Klassen *Acker-* und *Grünland* ist es wichtig, dass ein Pixel breite und geschlossene Linien an den Grenzen vorliegen. Geschlossene Linien sind notwendig um Feldsegmente voneinander abgrenzen zu können; die Grenzen sollen nur ein Pixel breit sein, damit möglichst viele Pixel pro Feldsegment zur Bestimmung der Merkmale für die Klassifizierung genutzt werden können. Im Folgenden soll daher nur auf Publikationen eingegangen werden, die dieses Kriterium erfüllen. Luo und Guo (2003) haben einen generellen Gruppierungsalgorithmus, basierend auf MRF, vorgestellt, der einzelne Segmenteigenschaften wie Fläche, Konvexität und Farbvarianz sowie paarweise Eigenschaften wie die Farbdifferenz und die Stärke einer Linie entlang der gemeinsamen Grenze für die Segmentierung verwendet. Dieser Algorithmus benötigt ein Training. Ein Training ist für die Anwendung der Verifikation von GIS-Ackerland- und Grünlandobjekten ungeeignet, da bedingt durch die Variation der unterschiedlichen Feldsegmente die benötigten Trainingsstichproben zu umfangreich sein würden.

Grote et al. (2007) benutzt mittlere Farbdifferenzen, die Stärke der gemeinsamen Grenze und die Standardabweichung der Farbe, um Regionen von Objekten iterativ zu verschmelzen, nachdem eine Übersegmentierung mittels des Normalized Cut Verfahrens berechnet wurde. Ein Training für dieses Verfahren ist nicht notwendig. Die Merkmale, die für das Verschmelzen der Regionen verwendet werden, sind hochgradig auf die Extraktion von Straßensegmenten spezialisiert. Es ist daher zu bezweifeln, dass diese Kriterien auch auf die Segmentierung von Schlägen in Ackerland- und Grünlandobjekten angewendet werden können. Zudem ist dieses Verfahren sehr laufzeitintensiv.

Ein weit verbreitetes Segmentierungsverfahren ist das Wasserscheidenverfahren, welches zur Segmentierung von Regionen homogener Grauwerte auf ein Gradientenbild angewendet wird, das zuvor zur Erzeugung stabiler Ergebnisse geglättet wird (Gonzalez und Woods, 2002). Zur Glättung dient in (Gonzalez und Woods, 2002) ein Gauß-Filter G_σ mit dem Glättungsparameter σ . Die Wasserscheidensegmentierung liefert Segmente mit geschlossenen und relativ glatten Grenzen, auch wenn es oft eine starke Übersegmentierung erzeugt (Tönnies, 2005). Um diesen Problemen zu begegnen, kann eine iterative Vorgehensweise gewählt werden. Zunächst wird eine Wasserscheidensegmentierung angewendet. Hierbei wird ein geringer Glättungsgrad angewendet, da eine starke Übersegmentierung verursacht werden soll. Anschließend wird ein Regionennachbarschaftsgraph (*region adjacency graph*, *RAG*) generiert, der aus Knoten und Kanten besteht. Die Knoten des Graphen sind die Regionen, während die Kanten die Nachbarschaftsrelationen repräsentieren. Knoten und Kanten besitzen Attribute, die Grundlage für die Verschmelzung ähnlicher Regionen sind. Dieser Ansatz wird u.a. in (Saarinen, 1994) verfolgt. Dieser Prozess wird durch wenige Parameter gesteuert, die einfach zu interpretieren sind. Ein Training des Verfahrens ist nicht erforderlich, und es liegen ein Pixel breite und geschlossene Linien an den Grenzen vor. Ein solches Segmentierungsverfahren erfüllt daher die oben genannten Anforderungen. Aus diesem Grund wird ein solches Verfahren in dieser Arbeit verwendet.

2.4. Diskussion

Beim Verifikationsprozess gibt es unterschiedliche Wege, GIS-Daten und das damit gegebene Vorwissen zu nutzen. In den meisten Fällen werden in der vorliegenden Literatur die Grenzen der GIS-Objekte als Basis für eine objektbasierte Klassifikation genutzt. Weiterhin werden die GIS-Daten immer für den finalen Schritt eines Verifikationsprozesses, dem Vergleich des Klassifikationsergebnisses der zu verifizierenden GIS-Objekt und der in der GIS-Datenbank eingetragenen Objektklasse, eingesetzt. Ein Verfahren, das erfolgreich zur Verifikation eines GIS eingesetzt wird, stellt Busch et al. (2005) vor. Bei dem Verfahren von Busch et al. (2005) handelt es sich um ein allgemeines und einfach übertragbares Verfahren, in dem neue Bildanalyseoperatoren zur Verifikation neuer Objektarten integriert werden können. Mit dem in (Busch et al., 2005) vorgestellten Verfahren ist zwar die Verifikation der gemeinsamen Objektklasse *Acker-/Grünland*, aber nicht eine getrennte Betrachtung dieser beiden Objektklassen möglich. Schwerpunkte der Arbeiten, die sich mit den Klassen *Acker-* und *Grünland* auseinandersetzen, sind nicht im Bereich der Verifikation, sondern im Bereich der Klassifikation zu finden.

Grundlage für eine erfolgreiche Verifikation von GIS-Objekten ist ein zuverlässiges Klassifikationsergebnis. Die Wahl des Klassifikationsverfahrens kann den Erfolg einer Klassifikation bei der Verwendung derselben Merkmale beeinflussen (Hoberg und Rottensteiner, 2010). Verfahren, die die gemeinsame Klasse *Acker-/Grünland* gut von anderen Klassen wie *Siedlung*, *Industrie* und *Wald* trennen kann, sind die Verfahren von Busch et al. (2005) sowie Rengers und Prinz (2009). In diesen beiden Verfahren wird gezeigt, dass eine zuverlässige Trennung der gemeinsamen Klasse *Acker-/Grünland* möglich ist, aber eine Trennung von *Acker-* und *Grünland* nicht erfolgreich durchgeführt werden kann. Während Rengers und Prinz (2009) nur textuelle Merkmale für die Klassifikation benutzen und wertvolle Farbinformationen bei der Klassifikation nicht berücksichtigen, verwendet Busch et al. (2005) zur Klassifikation das Verfahren von Gimel'farb (1996) unter Anwendung von MRF. Dabei werden auch die Farbinformationen genutzt. MRF sind gut geeignet für die Abtrennung von *Acker-/Grünland* zu anderen Klassen, da zum einen das Verfahren der MRF für mehrere Klassen geeignet ist und zum anderen lokales Kontextwissen eingebracht werden kann, was wie in Busch et al. (2005) gezeigt, zu einer zuverlässigen Trennung der Klassen *Acker-/Grünland*, *Siedlung*, *Industrie* und *Wald* beiträgt. Das Trennen von *Acker-* und *Grünland* würde dann ein Zwei-Klassen-Problem darstellen. Ein Klassifikationsverfahren, das in den unterschiedlichsten Bereichen der Wissenschaft und Technik Anwendung gefunden hat, ist das Verfahren der SVM. Ein wesentlicher Vorteil der SVM gegenüber anderen Klassifikationsverfahren ist, dass mit Klassen, die jeweils räumlich getrennte Clustern im Merkmalsraum besitzen, umgegangen werden kann. Räumlich getrennte Cluster im Merkmalsraum sind für die Klasse *Ackerland* auf Grund der unterschiedlichen Erscheinungsformen zu erwarten.

Ein anderer Aspekt der Klassifikation bezieht sich darauf, ob eine Klassifikation pixel- oder objektbasiert durchgeführt wird. In der Literatur wird der Begriff Objekt oft für Segmente benutzt, deren Pixel ähnliche Eigenschaften besitzen. Ein Vorteil der objektbasierten Klassifikation besteht darin, dass neben Eigenschaften eines Pixels auch Eigenschaften der Summe aller Pixel und die des Segmentes selber verwendet werden. Dadurch ist die objektbasierte Klassifikation sehr robust. Aber dadurch, dass nur das gesamte Segment einer Klasse zugeordnet wird, bleiben kleine Flächen anderer Nutzung innerhalb eines

Segmentes unberücksichtigt. Diese Flächen können durch eine pixelbasierte Klassifikation ermittelt werden, wobei eine Abtrennung dieser Flächen gegenüber Rauschen, z.B. verursacht durch den „Salt-and-Pepper-Effect“, schwierig ist. Eine Kombination eines pixel- und objektbasierten Verfahrens ist daher sinnvoll. Während mittels des pixelbasierten Verfahrens innerhalb eines GIS-Acker-/Grünlandobjektes Flächen mit anderer Nutzung detektiert werden können, verspricht eine objektbasierte Klassifikation eine bessere Möglichkeit der Trennung von Acker- und Grünland, da zusätzliche Merkmale (strukturelle und geometrische) für die Klassifikation verwendet werden können.

Egal ob die Klassifikation pixel- oder objektbasiert (also segmentbasiert) durchgeführt wird, das Ergebnis der Verifikation bezieht sich immer auf ein GIS-Objekt, d.h. die Ergebnisse einer pixel-/objektbasierten Klassifikation müssen immer auf ein GIS-Objekt übertragen werden, es sei denn, dass im Fall der objektbasierten Klassifikation ein Schlagkataster vorliegt. Das Problem der Segmentierung von GIS-Objekten in unterschiedliche Segmente, korrespondierend zu unterschiedlichen Bewirtschaftungseinheiten, wird in der Literatur i.d.R. umgangen, indem ein Schlagkataster verwendet wird, wie z.B. in den Arbeiten von Ruiz et al. (2004, 2007), Trias-Sanz (2006), Balaguer et al. (2010) sowie Recio et al. (2011). Ein Schlagkataster liegt in realen GIS-Datenbanken aber in den wenigsten Fällen vor und muss vom Benutzer erst hergestellt werden. Bei den Verfahren, die eine Segmentierung eines GIS-Objektes vornehmen, konnten entweder bei Klassifikationsverfahren keine zufriedenstellenden Ergebnisse erreicht werden (Peled und Gilichinsky, 2010) oder es werden simple Schachbrettsegmentierungen durchgeführt, bei denen die Segmente nicht mit den Grenzen der Schläge übereinstimmen (Rabatel et al., 2008), (Delenne et al., 2008), (Rengers und Prinz, 2009). Alternativ können die Segmente innerhalb eines GIS-Objektes mittels automatischer Segmentierung bestimmt werden. Als geeignetes Verfahren zur Segmentierung hat sich im Literaturüberblick ein Verfahren, basierend auf einer ursprünglich stark übersegmentierenden Wasserscheidentransformation mit anschließendem Verschmelzen von benachbarten Regionen mit ähnlichen Eigenschaften herausgestellt.

Neben der Wahl des richtigen Klassifikationsverfahrens hat nach Huang et al. (2000) auch die Auswahl geeigneter Merkmale großen Einfluss auf die Qualität des Klassifikationsergebnisses. Für die Klassifikation landwirtschaftlicher Flächen werden bisher Merkmale aus den Merkmalsgruppen der spektralen, textuellen, strukturellen und geometrischen Merkmale sowie weitere Merkmale verwendet. Rein spektrale Merkmale werden fast ausschließlich bei der multitemporalen Auswertung niedrig aufgelöster Satellitendaten genutzt (Itzerott und Kaden, 2006), (Itzerott und Kaden, 2007), (Marcal und Cunha, 2007), (Hoberg und Müller, 2011). Bei der Verwendung von textuellen Merkmalen kann zwar die gemeinsame Klasse *Acker-/Grünland* erfolgreich von anderen Klassen getrennt werden, eine Trennung der Klassen *Acker-* und *Grünland* ist jedoch nicht möglich (Busch et al., 2005), (Rengers und Prinz, 2009), obwohl der in (Busch et al., 2005) verwendete Ansatz von Gimel'farb (1996) sogar spektrale Eigenschaften berücksichtigt. Verfahren, die auf strukturellen Merkmalen basieren (Ranchin et al., 2001), (Wassenaar et al., 2002), (Chanussot et al., 2005), (Warner und Steinmaus, 2005), (Delenne et al., 2007), (Rabatel et al., 2008), (Delenne et al., 2008), benötigen häufig Vorinformationen, die bei Objektklassen wie *Weingarten* einfacher verfügbar und meist bekannt sind, jedoch nicht bei Ackerlandflächen, da bei Ackerflächen die Abstände der Pflanzreihen bzw. Bewirtschaftungsspuren je nach Fruchtart und verwendeter Bewirtschaftungsmaschinen stark schwanken können. Bei der monotemporalen Auswertung mit hochaufgelösten Satellitendaten müssen daher Merkmale aus mindestens zwei Merkmalsgruppen verwendet werden (Peled und Gilichinsky, 2010), (Haralick et al., 1973), (Gong et al., 2003), (Ranchin et al., 2001), (Ruiz et al., 2004), (Ruiz et al., 2007), (Balaguer et al., 2010), (Recio et al., 2011). Eine sichere Trennung der Klassen *Acker-* und *Grünland* wurde bisher nur erreicht, wenn Merkmale aus allen Merkmalsgruppen in einer objektbasierten Klassifikation verwendet werden (Ruiz et al., 2004), (Ruiz et al., 2007), (Balaguer et al., 2010), (Recio et al., 2011).

Zusammenfassend muss festgestellt werden: es gibt kein Verfahren zur Verifikation der getrennten Klassen *Acker-* und *Grünland*. Verfahren, die sich mit den getrennten Klassen *Acker-* und *Grünland* auseinandersetzen, beschränken sich auf die Klassifikation. Bei einer Klassifikation werden GIS-Spezifikationen wie z.B. das Vorhandensein mehrerer Schläge innerhalb eines Objektes von GIS mit dem Inhalt und Detaillierungsgrad einer topographischen Karte mittleren Maßstabs nicht berücksichtigt. Die Lösungen setzen ein Schlagkataster voraus oder bedienen sich einfacher Segmentierungsverfahren, bei denen die Segmentgrenzen nicht denen der Schläge entsprechen. Auch gibt es kein Verfahren, das bei der Klassifikation von Ackerland- und Grünlandobjekten eine Teilnutzung auf diesen GIS-Objekten detektiert und diese bei der Klassifikation berücksichtigt.

3. Grundlagen

Im diesem Kapitel sollen die Grundlagen für die in dieser Arbeit neu entwickelten Ansätze im Rahmen des Verifikationsverfahrens näher erläutert werden. Dazu gehört zum einen eine Einführung in den Klassifikationsalgorithmus der *Support Vector Machine* (Abschnitt 3.1) und zum anderen die Grundlagen für die Segmentierung eines GIS-Objektes (Abschnitt 3.2).

3.1. Einführung in den Klassifikationsalgorithmus der Support Vector Machine

Ein wesentlicher Bestandteil des Verifikationsprozesses ist die Klassifikation. Dazu werden pro Pixel/Objekt Merkmale berechnet, die in einem Merkmalsvektor gespeichert werden. Ziel der Klassifikation ist die Zuordnung eines Pixels/Objektes zu einer Objektklasse an Hand seines Merkmalsvektors. Das in dieser Arbeit verwendete Klassifikationsverfahren ist das der *Support Vector Machine* (SVM), ein überwachter Klassifikationsalgorithmus, eingeführt von Vapnik (1998). Die SVM stellt ein weit verbreitetes Klassifikationsverfahren dar und hat in den unterschiedlichsten Bereichen der Wissenschaft und Technik Anwendung gefunden, u.a. in der Fassadenrekonstruktion (Römer und Plümer, 2010), der Straßenextraktion (Fujimura et al., 2008), der Untersuchung von Hangrutschung (Heinert und Riedel, 2010), der Bestimmung tektonisch aktiver Gebiete (Herbert und Riedel, 2010), der Schrifterkennung (Schölkopf et al., 1996), der Veränderungsdetektion in Stadtgebieten (Nemmour und Chibani, 2006) und der Gesichtserkennung (Lu et al., 2001).

In diesem Abschnitt wird zunächst die Funktionsweise der SVM detailliert erklärt. Die Erläuterungen wurden vorwiegend aus Burges (1998), Bishop (2006) und Heinert (2010) entnommen. Anschließend werden wichtige Erweiterungen, Hinweise zur praktischen Verwendung einer SVM, die automatische Bestimmung der freien Parameter der SVM sowie das Lösen eines Mehrklassenproblems mittels SVM beschrieben. Abgeschlossen wird der Abschnitt mit einem Vergleich der SVM mit anderen Klassifikationsverfahren und einer kurzen Zusammenfassung.

3.1.1. Funktionsweise einer SVM

Ziel der Klassifikation ist die Zuordnung von Pixeln/Objekten an Hand von beobachteten Merkmalen zu einer von zwei Klassen (binäres Klassifikationsverfahren). Die beobachteten Merkmale eines Pixels/Objektes i , die z.B. aus einem Bild oder einem GIS abgeleitet werden können, werden in einem Merkmalsvektor³ $\mathbf{x}_i$ mit $\mathbf{x}_i \in \mathbb{R}^n$ gespeichert und sind die Eingangswerte in das Klassifikationsverfahren. n entspricht der Dimension des Merkmalsvektors, also der Anzahl der Merkmale. Das Ergebnis der Klassifikation ist die Klassenzugehörigkeit y der Pixel/Objekte mit $y \in \{-1; 1\}$, wobei $y = -1$ bzw. $y = +1$ die Zugehörigkeit zu jeweils einer der beiden Klassen (z.B. „Objekt“ und „Hintergrund“) beschreibt.

Da es sich um einen überwachten Klassifikationsalgorithmus handelt, ist ein Training notwendig. Trainingsdaten bestehen aus l Pixeln/Objekten mit bekannten Merkmalsvektoren $\mathbf{x}_i$ und bekannten Klassenzugehörigkeiten y_i mit $i \in \{1, \dots, l\}$. Die SVM trennt zwei Klassen an Hand ihrer Merkmalsvektoren durch eine Hyperebene H_d im Merkmalsraum so, dass der Abstand d der Hyperebene von denjenigen Vektoren, die der Hyperebene am nächsten liegen, maximiert wird (s. Abbildung 1). Dieses Prinzip wird Prinzip des maximalen Trennbereiches (*maximum-margin*) genannt.

Jene Vektoren, die genau diesen maximalen Abstand d haben, werden Stützvektoren (*Support Vectors*) genannt und sind in Abbildung 1 mit Kreisen gekennzeichnet. Die *Support Vectors* definieren zwei Ebenen H_1 und H_2 , welche parallel zu Hyperebene H_d liegen und von H_d den Abstand d haben. H_1 und H_2 definieren den Trennbereich (*margin*), innerhalb dessen sich keine Merkmalsvektoren befinden. Ziel des Trainings der SVM ist die Bestimmung der *Support Vectors*, die die Lage der Trennfläche im Merkmalsraum zur Bestimmung der Klassenzugehörigkeit eines unbekannten Pixels/Objektes bei gegebenen Merkmalsvektoren definieren.

³ kurz: Vektoren

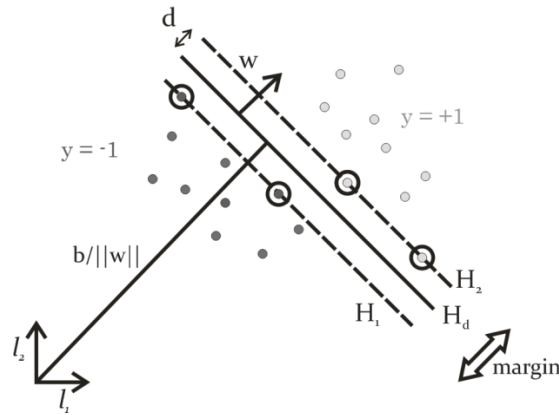


Abbildung 1: Hyperebene H_d für den linear trennbaren Fall zweier Klassen $y = -1$ und $y = +1$ (hell- und dunkelgraue Punkte), *Support Vectors* sind durch Kreise gekennzeichnet, alle anderen Symbole sind im Text erläutert.

Geht man zunächst von einem linear trennbaren Datensatz aus, gilt für auf der Hyperebene H_d liegende Vektoren $\mathbf{x}$ die Gleichung der Hyperebene $\mathbf{w} \cdot \mathbf{x} + b = 0$, wobei $\mathbf{w}$ die Normale zur Hyperebene, $b/\|\mathbf{w}\|$ die lotrechte Distanz der Hyperebene zum Koordinatenursprung und $\|\mathbf{w}\|$ die euklidische Norm von $\mathbf{w}$ ist. Für die *Support Vectors*, die auf den Ebenen H_1 und H_2 liegen, gilt

$$H_1, H_2: \pm d = \frac{\mathbf{w}'}{\|\mathbf{w}'\|} \cdot \mathbf{x} + \frac{b'}{\|\mathbf{w}'\|} \quad (3.1)$$

mit den Ebenenparametern $\mathbf{w}'$ und b' . Bei der Division von (3.1) durch den Abstand d ergibt sich

$$H_1, H_2: \pm 1 = \frac{\mathbf{w}'}{\|\mathbf{w}'\|d} \cdot \mathbf{x} + \frac{b'}{\|\mathbf{w}'\|d} \quad (3.2)$$

und unter der Annahme $\mathbf{w} = \mathbf{w}'/(\|\mathbf{w}'\|d)$ und $b = b'/(\|\mathbf{w}'\|d)$ gilt

$$H_1, H_2: \pm 1 = \mathbf{w} \cdot \mathbf{x} + b \quad (3.3)$$

In (3.3) gilt das positive Vorzeichen für Punkte mit $y_i = +1$ und das negative für $y_i = -1$. Anders ausgedrückt kann gesagt werden, dass

$$\mathbf{x}_i \cdot \mathbf{w} + b \geq +1 \quad \text{für } y_i = +1 \quad (3.4)$$

$$\mathbf{x}_i \cdot \mathbf{w} + b \leq -1 \quad \text{für } y_i = -1 \quad (3.5)$$

wobei für alle *Support Vectors* die Gleichheit und für alle anderen Vektoren die Ungleichheit gilt.

Aus (3.4) und (3.5) ergibt sich

$$y_i(\mathbf{x}_i \cdot \mathbf{w} + b - 1) \geq 0 \quad \forall i. \quad (3.6)$$

(3.6) muss für alle Trainingspunkte $(\mathbf{x}_i, y_i)$ erfüllt sein. Wenn *Support Vectors* in die Formeln (3.4), (3.5) und (3.6) eingesetzt werden, gilt die Gleichheit; in diesem Fall wird von einer aktiven Bedingung gesprochen, ansonsten von einer inaktiven Bedingung. Per Definition muss es immer mindestens zwei aktive Bedingungen geben, da mindestens zwei Punkte der Hyperebene am nächsten liegen (Bishop, 2006).

Wie bereits erwähnt, ist das Ziel, die Trainingsdaten nach dem Prinzip des *maximum-margin* voneinander zu trennen. Erreicht werden soll also, dass d maximiert wird. Das folgende Konzept der Maximierung von d wurde (Bishop, 2006) entnommen. Für d gilt die Gleichung $d = 1/\|\mathbf{w}\|$, was wie folgt gezeigt wird. Die Normalabstände der Ebenen H_1 und H_2 vom Ursprung ergeben sich aus:

$$H_1: \frac{\mathbf{x}_i \mathbf{w} + b + 1}{\|\mathbf{w}\|} = 0, \quad d_{origin}(H_1) = \frac{b+1}{\|\mathbf{w}\|} \quad (3.7)$$

und

$$H_2: \frac{\mathbf{x}_i \mathbf{w} + b - 1}{\|\mathbf{w}\|} = 0, \quad d_{origin}(H_2) = \frac{b-1}{\|\mathbf{w}\|}. \quad (3.8)$$

Da die Ebenen H_1 und H_2 parallel liegen, ergibt die Differenz Δ der Normalabstände den *margin*, welcher $2d$ entspricht. Daher gilt

$$\Delta = 2d = \frac{b+1}{\|\mathbf{w}\|} - \frac{b-1}{\|\mathbf{w}\|} = \frac{2}{\|\mathbf{w}\|} \quad (3.9)$$

und damit gilt $d = 1/\|\mathbf{w}\|$. Daher ist die Maximierung von d gleichbedeutend mit der Maximierung von $1/\|\mathbf{w}\|$. Dies ist wiederum äquivalent der Minimierung von $\|\mathbf{w}\|$ bzw. zur leichteren mathematischen Handhabung der Minimierung von

$$z = \frac{1}{2} \|\mathbf{w}\|^2 \quad (3.10)$$

mit den l Nebenbedingungen $y_i(\mathbf{x}_i \cdot \mathbf{w} + b - 1) \geq 0 \forall i$. Der Faktor $\frac{1}{2}$ in (3.10) wird aus Gründen der Vereinfachung für spätere Berechnungen eingeführt.

Ziel ist es, die quadratische Funktion (3.10) unter den gegebenen Nebenbedingungen zu minimieren. Es liegt ein Beispiel für ein quadratisches Optimierungsproblem mit Ungleichungen als Nebenbedingungen vor (Bishop, 2006). Es fällt auf, dass der Parameter b aus der Zielfunktion (3.10) verschwunden ist, b fließt jedoch in die Berechnung auf Grund der Nebenbedingungen implizit ein.

Um das Optimierungsproblem mit Nebenbedingungen zu lösen, werden Lagrangemultiplikatoren $\alpha_i, i = 1, \dots, l$ eingeführt, mit einem Lagrangemultiplikator $\alpha_i \geq 0$ für jede Nebenbedingung. Die Lagrangefunktion L lautet dann

$$L(\mathbf{w}, b, \alpha) \equiv \frac{1}{2} \|\mathbf{w}\|^2 - \sum_{i=1}^l \alpha_i y_i (\mathbf{x}_i \cdot \mathbf{w} + b) + \sum_{i=1}^l \alpha_i y_i \rightarrow \min \quad (3.11)$$

mit den Nebenbedingungen $\alpha_i \geq 0 \forall i$. Die Lagrangefunktion L wird minimiert in Hinsicht auf $\mathbf{w}$ und b und maximiert in Hinsicht auf α (Bishop, 2006). Setzt man die Ableitungen von L nach $\mathbf{w}$ und b zu Null, ergeben sich die Bedingungen

$$\frac{\partial L(\mathbf{w}, b, \alpha)}{\partial \mathbf{w}} = \frac{1}{2} \cdot 2\mathbf{w} - \sum_{i=0}^l \alpha_i y_i \mathbf{x}_i \stackrel{!}{=} 0 \quad (3.12)$$

oder

$$\mathbf{w} = \sum_{i=0}^l \alpha_i y_i \mathbf{x}_i \quad (3.13)$$

sowie

$$\frac{\partial L(\mathbf{w}, b, \alpha)}{\partial b} = - \sum_{i=0}^l \alpha_i y_i \stackrel{!}{=} 0 \quad (3.14)$$

oder

$$\sum_{i=0}^l \alpha_i y_i = 0. \quad (3.15)$$

Werden die Gleichungen (3.13) und (3.15) in die Gleichung (3.11) eingesetzt, ergibt sich nach Umformung (Bishop, 2006):

$$L_D(\alpha) = \sum_i \alpha_i - \frac{1}{2} \sum_{ij} \alpha_i \alpha_j y_i y_j \mathbf{x}_i \cdot \mathbf{x}_j \rightarrow \max \quad (3.16)$$

mit den Nebenbedingungen $\alpha_i \geq 0$ und $\sum \alpha_i y_i = 0$, wobei nur für die *Support Vectors* $\alpha_i > 0$ gilt. Für alle anderen Punkte gilt $\alpha_i = 0$. L_D ist die duale Formulierung von L , d.h. eine Lösung des Problems wird durch die Minimierung von L (3.11) bzw. durch die Maximierung von L_D (3.16) erreicht (Bishop, 2006).

Bei (3.16) handelt es sich um eine quadratische Zielfunktion mit linearen Bedingungen, so dass die Lagrangemultiplikatoren α_i mittels quadratischer Optimierung bestimmt werden können, wie in (Bishop, 2006) dargestellt. Nachdem dies erfolgt ist, können $\mathbf{w}$ und b mittels der *Support Vectors* mit $\alpha_i > 0$ berechnet werden. Zunächst erfolgt die Berechnung von $\mathbf{w}$ mittels (3.13). Anschließend kann mittels (3.3) für jeden *Support Vector* ein Wert für b bestimmt werden. Diese Werte für b werden anschließend gemittelt.

Damit sind alle Unbekannten aus Gleichung (3.6) bestimmt und die Klassifikation neuer Objekte mittels deren Merkmalsvektoren $\mathbf{x}$ kann durch Einsetzen in die Entscheidungsfunktion

$$y = \text{sign}(\sum_{i=1}^s \alpha_i y_i \mathbf{x}_i \cdot \mathbf{x} + b) \quad (3.17)$$

berechnet werden, wobei s die Anzahl der *Support Vectors* ist und die Summe nur über die *Support Vectors* läuft. Je nachdem, welches Vorzeichen sich ergibt, gehört das Objekt zu der Klasse -1 oder der Klasse $+1$. Jeder Merkmalsvektor mit $\alpha_i = 0$ aus (3.17) hat auf die Summe in (3.17) und damit auf die Klassenbestimmung keinen Einfluss. Lediglich die *Support Vectors* bestimmen die Lage der Trennfläche im Merkmalsraum und damit die Klassenzugehörigkeit eines neuen Objektes (3.17).

Bisher wurde nur der linear trennbare Klassifikationsfall, der in der Praxis selten auftritt, betrachtet. Unter anderen in (Heinert, 2010) wird beschrieben, wie mit dem nicht linear trennbaren Klassifikationsfall umgegangen werden kann. Danach wird eine Klassifikation von nicht linear trennbaren Merkmalsvektoren möglich, indem die Merkmalsvektoren vom Merkmalsraum in einen höher dimensional Raum transformiert werden, in dem eine lineare Trennbarkeit der transformierten Merkmalsvektoren erreicht werden kann (s. Abbildung 2).

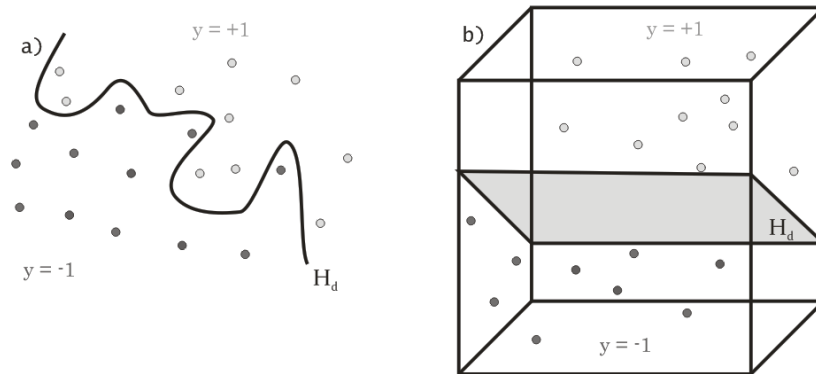


Abbildung 2: a) Nicht linear trennbarer Datensatz im Merkmalsraum b) Lineare Trennbarkeit des Datensatzes nach der Transformation in einen höher dimensional Raum.

Da es sehr aufwendig sein kann, große Mengen von Merkmalsvektoren in einen höher dimensional Merkmalsraum zu transformieren, wird der sogenannte Kernel-Trick angewendet. Die Merkmalsvektoren erscheinen im Trainings- (3.16) und Klassifikationsproblem (3.17) nur in Form ihres Skalarproduktes $\mathbf{x}_i \cdot \mathbf{x}_j$ bzw. $\mathbf{x}_i \cdot \mathbf{x}$. Wenn die Vektoren des Merkmalsraums in einen höher dimensional Merkmalsraum mit Hilfe einer Mappingfunktion Φ transformiert werden, dann liegen anstelle der Skalarprodukte die Skalarprodukte der Mappingfunktionen vor, also

$$L_D(\alpha) = \sum_i \alpha_i - \frac{1}{2} \sum_{ij} \alpha_i \alpha_j y_i y_j \Phi(\mathbf{x}_i) \cdot \Phi(\mathbf{x}_j) \rightarrow \max, \text{ und} \quad (3.18)$$

$$f(\mathbf{x}) = \text{sign}(\sum_{i=1}^S \alpha_i y_i \Phi(\mathbf{x}) \cdot \Phi(\mathbf{x}_i) + b) \quad (3.19)$$

Wenn nun eine Kernfunktion K mit $K(\mathbf{x}_i, \mathbf{x}_j) = \Phi(\mathbf{x}_i) \cdot \Phi(\mathbf{x}_j)$ existiert, dann kann diese verwendet werden, ohne dass Φ explizit bestimmt werden muss. Eine solche Kernfunktion ist stetig symmetrisch und muss positiv definit sein. Positiv definit ist die Kernfunktion genau dann, wenn die Mercer-Bedingung erfüllt ist. Die Mercer-Bedingung beschreibt nicht wie die Kernfunktion gebildet werden kann, sondern nur ob eine Funktion eine geeignete Kernfunktion sein kann. Nähere Details sind (Burgess, 1998) und (Heinert, 2010) zu entnehmen.

Vorteile der Verwendung einer Kernfunktion sind in (Bousquet und Schölkopf, 2006) beschrieben. Ein Vorteil der Verwendung einer Kernfunktion ist die einfache Handhabung von nicht linear trennbaren Klassifikationsfällen, oft ohne zusätzlichen rechnerischen Aufwand, da sie eine explizite Transformation in den höher dimensional Merkmalsraum vermeidet. Als weiterer Vorteil wird die Konvexität des dazugehörigen Optimierungsproblems genannt.

Die Wahl der Kernfunktion entscheidet über die Zuverlässigkeit des Ergebnisses. Die Beschränkung auf **eine** Kernfunktion ist daher als Problem zu sehen. Nachdem die Entscheidung auf einen Kern gefallen ist, kann der Benutzer die Berechnung nur durch die Parameter der Kernfunktion beeinflussen. Der Nachteil der Beschränkung auf eine Kernfunktion liegt darin, vorab die „richtige“ Wahl der Kernfunktion für eine definierte Aufgabe treffen zu müssen. Die „richtige“ Wahl der Kernfunktion für eine definierte Aufgabe ist immer noch ein Forschungsgebiet. Eine Möglichkeit, diesem Nachteil entgegenzutreten ist es, verschiedene Kernfunktionen für das Training zu nutzen (Bach et al., 2004).

Eine häufig verwendete Kernfunktion ist die Gauß-Kernfunktion (oder auch radiale Basisfunktion kurz RBF-Kernfunktion genannt) mit

$$K(\mathbf{x}_i, \mathbf{x}_j) = e^{(-\gamma \|\mathbf{x}_i - \mathbf{x}_j\|^2)}, \quad (3.20)$$

wobei $2/\gamma$ die Breite der Gaußfunktion ist. Ein Vorteil der RBF-Kernfunktion ist, dass die Kernfunktion von nur einem Parameter γ abhängig ist, während die meisten anderen Kernfunktionen mehrere Parameter benötigen. Bei Vergleichen der RBF-Kernfunktion mit anderen Kernfunktionen, z.B. mit einer linearen Kernfunktion (Römer und Plümer, 2010), ANOVA-Kernfunktion, neuronalen Kernfunktion (Heinert und Riedel, 2010) oder polynomialen Kernfunktion (Huang et al., 2000), erzielte die RBF-Kernfunktion bei jedem Test das beste Ergebnis bzw. genauso gute Ergebnisse wie die jeweils beste andere Kernfunktion.

Die Verwendung einer Kernfunktion führt ausgehend von (3.17) bzw. (3.18) zu der Entscheidungsfunktion für die Klassifikation von

$$f(\mathbf{x}) = \text{sign}(\sum_{i=1}^S \alpha_i y_i K(\mathbf{x}_i, \mathbf{x}) + b). \quad (3.21)$$

Analog zu (3.17) bzw. (3.18) wird auch hier die Summe nur über die *Support Vectors* berechnet.

3.1.2. SVM mit Fehlern in den Trainingsdaten

Der bisher vorgestellte Lösungsansatz gilt nur für sich nicht überschneidende Klassen im Merkmalsraum. Sich überschneidende Klassen werden auch überlappende Klassen genannt. Handelt es sich um Daten, bei denen sich die Trainingsdaten im Merkmalsraum überlappen, führt das Lernen zwar in der Regel zu einer genauen Trennung der Vektoren im Merkmalsraum. Dies kann aber zu einer Überanpassung an den Trainingsdatensatz führen, wie im Beispiel in Abbildung 3a) zu sehen ist. Aus diesem Grund muss der SVM erlaubt werden, einige Trainingsobjekte falsch klassifizieren zu können. Es wird dann erlaubt, dass Trainingsvektoren auf der falschen Seite der Ebenen H_d liegen können (s. Abbildung 3b). Dadurch wird die SVM weniger anfällig gegenüber der Überanpassung an die Trainingsdaten (Bousquet und Schölkopf, 2006).

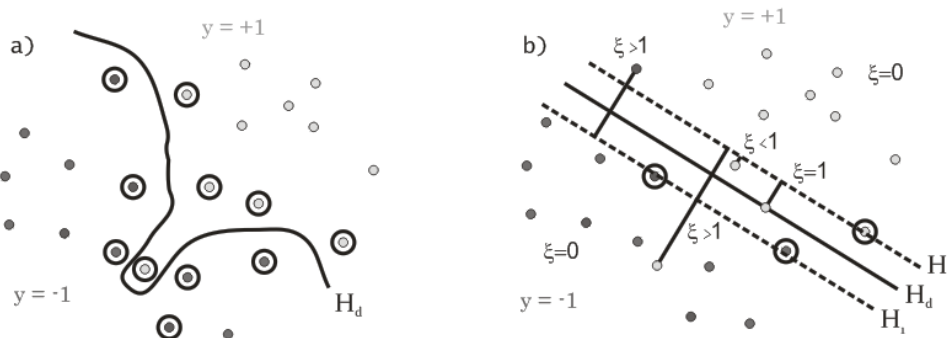


Abbildung 3: a) Überanpassung der SVM an sich überlappende Trainingsdaten, b) Einführung von Schlupfvariablen ξ_i (Support Vectors sind durch Kreise gekennzeichnet).

Dies wird mit der Einführung von nicht negativen Schlupfvariablen (*slack variables*) $\xi_i, i = 1, \dots, l$ mit $\xi_i \geq 0$ für jeden Vektor erreicht (Heinert, 2010). Bei Einführung der Schlupfvariablen in (3.4) und (3.5) ergibt sich

$$(\mathbf{x}_i \cdot \mathbf{w} + b) \geq +1 - \xi_i \quad \text{für } y_i = +1 \quad (3.22)$$

$$(\mathbf{x}_i \cdot \mathbf{w} + b) \leq -1 + \xi_i \quad \text{für } y_i = -1 \quad (3.23)$$

mit $\xi_i \geq 0$. Aus (3.22) und (3.23) ergibt sich analog zu (3.6)

$$y_i(\mathbf{x}_i \cdot \mathbf{w} + b) \geq 1 - \xi_i \quad \forall i = 1 \dots N \quad (3.24)$$

Liegt ein Vektor $\mathbf{x}_i$ auf der Ebene H_1 bzw. H_2 oder auf der Seite der richtigen Klasse, gilt $\xi_i = 0$. Befindet sich ein Vektor $\mathbf{x}_i$ auf der Hyperebene H_d , dann beträgt $\xi_i = 1$. Vektoren, für die $0 < \xi_i \leq 1$ gelten, liegen innerhalb des margin, aber auf der richtigen Seite der Hyperebene H_d . Vektoren mit $\xi_i > 1$ liegen auf der falschen Seite der Hyperebene H_d (s. Abbildung 3b).

Unser Ziel ist es nun, die Hyperebene so zu bestimmen, dass der *margin* maximal wird und gleichzeitig möglichst wenige Vektoren, die auf der falschen Seite der Ebenen H_d oder im *margin* liegen, zuzulassen. Dies erreicht man durch Minimierung von

$$\frac{\|\mathbf{w}\|^2}{2} + C(\sum_i \xi_i) \rightarrow \min. \quad (3.25)$$

C ist ein Strafterm, der die „Kosten“ für einen Punkt im *margin* oder jenseits der Hyperebene mit $\xi_i > 0$ beschreibt. Für ein sehr großes C entsteht eine Hyperebene wie in Abbildung 3a) gezeigt.

Die Minimierung von (3.25) unter Berücksichtigung der Nebenbedingungen (3.24) der weiteren Nebenbedingung $\xi_i \geq 0$ führt zu der Lagrangefunktion

$$L(\mathbf{w}, b, \alpha) \equiv \frac{1}{2} \|\mathbf{w}\|^2 + C \sum_{i=1}^l \xi_i - \sum_{i=1}^l \alpha_i \{y_i(\mathbf{x}_i \cdot \mathbf{w} + b) - 1 + \xi_i\} - \sum_{i=1}^l \mu_i \xi_i \rightarrow \min, \quad (3.26)$$

wobei α_i und μ_i Lagrangemultiplikatoren sind und die Nebenbedingungen

$$\alpha_i \geq 0, \quad (3.27)$$

$$\mu_i \geq 0, \quad (3.28)$$

$$\xi_i \geq 0 \quad (3.29)$$

$$\mu_i \xi_i = 0 \quad (3.30)$$

$$y_i(\mathbf{x}_i \cdot \mathbf{w} + b) - 1 + \xi_i \geq 0, \quad (3.31)$$

$$\alpha_i \{y_i(\mathbf{x}_i \cdot \mathbf{w} + b) - 1 + \xi_i\} = 0 \quad (3.32)$$

gelten. Die Gleichungen (3.30) und (3.32) können benutzt werden, um für b einen Grenzwert zu ermitteln. Nun wird (3.26) in Bezug auf $\mathbf{w}$, b und $\{\xi_i\}$ optimiert. Dazu werden die Ableitungen von L nach $\mathbf{w}$, b und ξ_i zu Null gesetzt. Daraus ergibt sich:

$$\frac{\partial L(\mathbf{w}, b, \alpha)}{\partial \mathbf{w}} = \frac{1}{2} \cdot 2\mathbf{w} - \sum_{i=0}^l \alpha_i y_i \mathbf{x}_i = 0, \quad (3.33)$$

das heißt, dass

$$\mathbf{w} = \sum_{i=0}^l \alpha_i y_i \mathbf{x}_i \quad (3.34)$$

gilt. Des Weiteren lauten die Ableitungen

$$\frac{\partial L(\mathbf{w}, b, \alpha)}{\partial b} = - \sum_{i=0}^l \alpha_i y_i = 0 \quad (3.35)$$

und

$$\frac{\partial L(\mathbf{w}, b, \alpha)}{\partial \xi_i} = C - \alpha_i - \mu_i = 0, \quad (3.36)$$

das heißt, dass

$$\mu_i = C - \alpha_i \quad (3.37)$$

gilt. Werden nun (3.34), (3.35) und (3.37) in (3.26) eingesetzt, erhält man

$$L_D(\alpha) = \sum_i \alpha_i - \frac{1}{2} \sum_i \alpha_i \alpha_j y_i y_j \mathbf{x}_i \cdot \mathbf{x}_j \rightarrow \max, \quad (3.38)$$

das der dualen Funktion von (3.16) mit den Nebenbedingungen $0 \leq \alpha_i \leq C \forall i = 1, \dots, l$ und $\sum \alpha_i y_i = 0$ entspricht. Wie zu sehen, taucht ξ in dieser Optimierung nicht mehr explizit auf. Die Lösung von (3.38) erfolgt analog der Lösung von (3.16).

Es sei darauf hingewiesen, dass Gleichung (3.37) kombiniert mit Gleichung (3.30) zeigt, dass $\xi_i = 0$ ist, wenn $\alpha_i < C$. Daher kann jeder Trainingspunkt, für den $0 < \alpha_i < C$ mit Gleichung (3.32) (mit $\xi_i = 0$) verwendet werden, um b zu berechnen.

3.1.3. Anmerkungen zur praktischen Durchführung

Ziel des Trainings der SVM ist die Bestimmung der *Support Vectors*, die die Lage der Trennfläche im Merkmalsraum zur Bestimmung der Klassenzugehörigkeit eines unbekannten Objektes bei gegebenem Merkmalsvektor definieren. Um eine Klassifikation erfolgreich durchführen zu können, müssen die

Trainingsdaten repräsentativ gewählt werden. In Untersuchungen von Brank et al. (2002) wurde festgestellt, dass die Auswahl der Trainingsgebiete, hier vor allem die Anzahl der Trainingsgebiete in dem Sinne, dass keine Klasse unterrepräsentiert ist, einen größeren Einfluss auf eine erfolgreiche Klassifikation hat als das Ermitteln relevanter Merkmale für die Klassifikation oder die Gewichtung/Selektion dieser Merkmale.

Bevor das Anlernen und auch später die Klassifikation jedoch durchgeführt werden können, sollten die Merkmale skaliert werden. Liegen i Merkmalsvektoren mit $\mathbf{x}_i = (x_1, \dots, x_k)^T$ vor, wobei x_j mit $j \in \{1, \dots, k\}$ ein Merkmal eines Merkmalsvektors ist, so werden die jeweiligen x_j aller Merkmalsvektoren skaliert. Die Skalierung erfolgt z.B. linear in den Bereich von -1 und +1, so dass $x_j \in [-1; 1]$. Es ist dabei für Training und Klassifikation dieselbe Skalierung zu verwenden. Der größte Vorteil der Skalierung liegt in der Vermeidung des Dominierens großer numerischer Werte gegenüber kleinen numerischen Werten (Hsu et al., 2010).

3.1.4. Automatische Bestimmung der Parameter der SVM

Parameter, die bei der SVM vom Benutzer gewählt werden müssen, sind der Parameter C (Strafwert für Fehlklassifikationen, siehe (3.25)) sowie die Parameter der jeweiligen Kernfunktion. Wird z.B. die RBF-Kernfunktion verwendet, handelt es sich um den Parameter γ (siehe (3.20)). Foody (2001) stellt in seinen Untersuchungen fest, dass bei einem zu großen Wert für γ und/oder einem zu großen Wert für C die Gefahr der Überanpassung (*overfitting*) an die Trainingsdaten besteht. Zu kleine Werte für γ und C führen hingegen zu einer zu starken Glättung und zu vielen Trainingsfehlern. Ziel ist es daher, jene Werte für die Parameter C und γ zu identifizieren, für die das beste Klassifizierungsergebnis erreicht werden kann. Dies kann mit Hilfe der Kreuzvalidierung (*cross-validation*) und der Gittersuche (*grid-search*) erreicht werden.

Bei der Kreuzvalidierung werden die Trainingsdaten in Untergruppen gleicher Größe aufgeteilt. Alle bis auf eine Untergruppe werden genutzt, um die SVM zu trainieren. Die so angelernete SVM wird an der verbleibenden Untergruppe validiert. Das Ergebnis ist die Prozentzahl der richtig klassifizierten Objekte. Im Weiteren wird der Vorgang mit anderen Untergruppen für Training und Test wiederholt.

Die Gittersuche muss mit der Kreuzvalidierung ausgeführt werden. Bei der Gittersuche werden systematisch verschiedene Parameterpaare für C und γ für eine Kreuzvalidierung verwendet (s. Abbildung 4). In der Implementierung von Hsu et al. (2010) wird der Parameter C mit $C = 2^{-5}, 2^{-3}, \dots, 2^{+15}$ und der Parameter γ mit $\gamma = 2^{-15}, 2^{-13}, \dots, 2^3$ getestet.

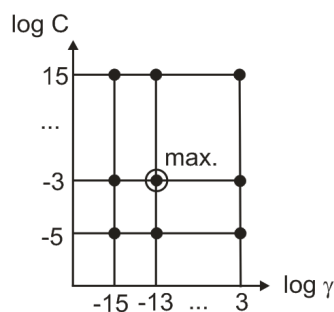


Abbildung 4: Schematische Darstellung der Gittersuche (ausgefüllte Punkte sind mögliche Parameterpaare, der Extrakreis kennzeichnet das Parameterpaar mit dem besten Kreuzvalidierungsergebnis).

Das Parameterpaar C und γ , für das das beste Kreuzvalidierungsergebnis erzielt werden konnte, also die höchste Prozentzahl der richtig klassifizierten Objekte, wird nach Abschluss der Suche aller möglichen Kombinationen aus der Gittersuche ausgewählt. Nach Hsu et al. (2010) kann die Kreuzvalidierung mit Gittersuche auf diesem Wege eine Überanpassung der SVM an die Trainingsdaten verhindern.

Für die angesprochene Gittersuche erfolgt die Auswahl der sich ändernden Parameter C und γ i.d.R. zunächst entsprechend einer logarithmischen Skala wie oben beschrieben. Außerdem wird die Strategie vom Groben ins Feine verwendet, d.h. mittels einer groben Suche wird zunächst ein Bereich für die Parameter C und γ eingegrenzt, z.B. der in Abbildung 4 mit einem Kreis gekennzeichnete. Anschließend wird dieser Bereich mit einer feineren Suche weiter abgetastet. Nähere Angaben sind Hsu et al. (2010) zu entnehmen.

Ein Nachteil der Kreuzvalidierung mit Gittersuche ist, dass dieser Prozess sehr rechenintensiv ist. Jedoch kann das Finden des optimalen Parameterpaares C und γ parallelisiert werden, da jedes Parameterpaar unabhängig voneinander ist und die Kreuzvalidierung deshalb unabhängig voneinander durchgeführt werden kann.

3.1.5. Anwenden einer SVM auf ein Mehrklassenproblem

Die SVM ist grundsätzlich ein Zwei-Klassen-Klassifikator. In der Praxis ist man jedoch häufig mit $K > 2$ Klassen konfrontiert. Die Behandlung eines Mehrklassenproblems mittels einer SVM ist durch die Kombination verschiedener Zwei-Klassen-SVM-Klassifikationen möglich. Folgende Ansätze sind Bishop (2006) entnommen.

Ein weit verbreiteter Ansatz ist die Benutzung multipler Zwei-Klassen-SVM-Klassifikatoren. Dabei werden K verschiedene SVM trainiert. Es wird eine der K Klassen gegen alle anderen $K - 1$ Klassen, die zu einer zweiten Klasse vereint werden, klassifiziert (*one-versus-all*). Für jede der K Klassen wird, z.B. der Abstand dieser Klasse von der Hyperebene gespeichert. Die Klasse, die am weitesten von der Hyperebene entfernt liegt, legt die zu bestimmende Klassenzugehörigkeit des Vektors fest (*Winner-Takes-All-Strategie*). Alternative Maße zur Bestimmung der Klassenzugehörigkeit eines Merkmalsvektors werden u.a. in (Kreßel, 2002) und (Liu und Zheng, 2005) beschrieben. Probleme bei diesem Vorgehen bestehen darin, dass jede der SVM auf unterschiedlichen Daten trainiert wird. Hinzu kommt, dass die Trainingsdatensätze unausgeglichen sind. So wird z.B. bei einem 10-Klassen-Problem eine Klasse mit 10% aller Trainingsdaten trainiert, während die zweite Klasse, die alle anderen Klassen vereint, 90% der Daten besitzt.

Ein anderer weit verbreiteter Ansatz eines multiplen Zwei-Klassen-Klassifikators besteht darin, $K(K - 1)/2$ Zwei-Klassen-SVM zu trainieren, um alle möglichen Kombinationen von Klassen gegeneinander zu testen (*one-versus-one*). Die Klasse, die die meisten Zuschläge erhält, wird als Gewinnerklasse ausgewählt und bestimmt die Klassenzugehörigkeit des Vektors (*Max-Wins-Voting*). Nachteil dieser Strategie gegenüber der oben vorgestellten *one-versus-all*-Lösung ist die wesentlich längere Berechnungszeit, da mehr Zwei-Klassen-SVM zu berechnen sind. Vorteil ist, dass der Algorithmus immer mit den gleichen Trainingsdaten pro Klasse trainiert wird.

3.1.6. Vergleich mit anderen Klassifikationsalgorithmen

Wie bereits vorgestellt, ist das ML-Verfahren ein gängiges Klassifikationsverfahren in der Fernerkundung. Bei dem ML-Verfahren wird davon ausgegangen, dass die Merkmalsvektoren eine bestimmte Verteilung besitzen. Wird diese Verteilungsannahme verletzt, kann es sein, dass der ML-Klassifikator inkonsistent wird. Der größte Vorteil der SVM gegenüber der ML-Klassifikation liegt daher darin, dass eine Annahme über die Verteilung der Merkmalsvektoren bei der Verwendung einer SVM nicht notwendig ist. Ein Vorteil der ML-Klassifikation gegenüber der SVM ist der bessere Umgang mit Ausreißern. Ausreißer sind Merkmalsvektoren, die fernab aller anderen Vektoren im Merkmalsraum liegen. Bei der SVM können Ausreißer explizit nicht erkannt werden, da nur die Lage der Vektoren gegenüber der Trennfläche, nicht aber die Entfernung zu allen anderen Vektoren für die Klassifizierung relevant ist. Das Erkennen von Ausreißern bei der Verwendung einer SVM wird u.a. in (Ziems et al., 2011) behandelt. Ein direkter Vergleich des ML-Klassifikators mit einer SVM (Mehrklassen-SVM mit RBF-Kernfunktion) fand in (Huang et al., 2000), (Foody und Mathur, 2004) und (Braun et al., 2010) statt. Bei allen Veröffentlichungen war die SVM signifikant besser als der ML-Klassifikator.

Ein weiteres ebenfalls bereits vorgestelltes Verfahren ist das kNN-Verfahren. Dieses Verfahren besitzt gegenüber der SVM den Nachteil, dass es unterschiedliche Verteilungen von Vektoren im Merkmalsraum nicht berücksichtigt. Beide Klassifikationsverfahren wurden u.a. in Schölkopf et al. (1996) verglichen. Bei dem Test schnitt die SVM deutlich besser ab.

Weitere Verfahren sind die Neuronalen Netze und die baumförmigen Klassifikationsverfahren (*Decision Tree*). Wie bereits vorgestellt, sind baumförmige Klassifikationsverfahren, u.a. (Quinlan, 1993), sehr flexibel, da für jede Einzelentscheidung die günstigste Kombination der Merkmale verwendet wird, sie neigen dadurch aber zur Überanpassung an die Trainingsdaten. Gerade die Strategie zur Vermeidung der Überanpassung an die Trainingsdaten, nicht nur gegenüber den baumförmigen Klassifikationsverfahren, ist ein Vorteil der SVM (Heinert, 2010). Die Trennfläche verfügt bei richtiger Anwendung der SVM über eine endliche und im günstigsten Fall eine signifikant kleinere Anzahl an *Support Vectors* gegenüber der Anzahl aller verfügbaren Vektoren. Alle anderen Vektoren, außer den *Support Vectors*, haben keinen Einfluss auf die Lage der Trennfläche und damit auf die Bestimmung der Klassenzugehörigkeit neuer Objekte. Sowohl die Neuronalen Netze als auch ein baumförmiges Klassifikationsverfahren wurden in (Huang et al., 2000) und (Foody und Mathur, 2004) mit einer SVM bezüglich der Klassifikation von Landnutzungsklassen verglichen. Die SVM war signifikant besser als die beiden anderen Verfahren.

Ein häufig genannter Vorteil der SVM gegenüber vielen anderen Klassifikationsverfahren ist die Eigenschaft der SVM, gute Genauigkeiten bei einer geringen Anzahl von Trainingsgebieten zu erreichen. Ein Vergleich unter diesem Gesichtspunkt mit Neuronalen Netzen, dem baumförmigen Klassifikationsverfahren nach Quinlan (1993) sowie dem ML-Klassifikationsverfahren ist in (Huang et al., 2000) gegeben. Bei der Nutzung aller zur Verfügung stehenden Merkmale für die Klassifikation erzielt die SVM immer das signifikant beste Ergebnis bei verschiedenen großen Trainingsgebieten.

Ein in jüngerer Zeit häufig benutztes Klassifikationsverfahren ist der ebenfalls bereits vorgestellte CRF-Klassifikator. In (Hoberg und Müller, 2011) wird die SVM mit dem CRF-Klassifikationsverfahren verglichen. Ergebnisse zeigen, dass der CRF-Klassifikator geringfügig bessere Ergebnisse als der SVM-Klassifikator erreicht. Das liegt vor allem daran, dass an den Grenzen zwischen Flächen verschiedener Nutzung durch die glättende Eigenschaft der CRF das Ergebnis des CRF-Klassifikators bei der in dieser Veröffentlichung verwendeten pixelbasierten Klassifikation besser geeignet ist. Werden nur die inneren Bereiche der Flächen betrachtet, sind die Ergebnisse beider Klassifikatoren vergleichbar.

Einer der größten Nachteile der SVM gegenüber allen zuvor genannten Klassifikationsalgorithmen liegt darin, dass es sich um ein binäres Klassifikationsverfahren handelt. Mehrklassenprobleme können bei der SVM nur über multiple Zwei-Klassen-SVMs gelöst werden, wobei alle anderen zuvor genannten Verfahren leicht an ein Mehrklassenproblem angepasst werden können. Ein weiterer Nachteil der SVM ist, dass kein direktes Qualitätsmaß abgeleitet werden kann. Nur über Umwege ist es möglich, ein solches Qualitätsmaß zu erhalten (Platt, 2000).

Des Weiteren ist das Lernen einer SVM sehr rechenintensiv, z.B. verglichen mit dem ML-Klassifikator. Der Lernprozess der SVM wird durch das zusätzliche Lernen der Parameter C sowie der Kernfunktionsparameter mittels Gittersuche und Kreuzvalidierung weiter verlangsamt. Werden Mehrklassenprobleme angelernt, verlangsamt sich der Prozess zusätzlich.

3.1.7. Zusammenfassung

In diesem Abschnitt wurde zunächst die theoretische Grundlage der SVM im Detail erörtert. Der binäre überwachte Klassifikationsalgorithmus der SVM trennt zwei Klassen mittels ihrer Merkmalsvektoren durch eine Hyperebene so, dass der Abstand derjenigen Vektoren, die der Hyperebene am nächsten liegen, maximiert wird. Die SVM arbeitet sowohl auf im Merkmalsraum linear als auch auf nicht linear trennbaren Daten. Das wird durch eine Transformation der nicht linear trennbaren Daten des Merkmalsraums in einen höher dimensional Merkmalsraum erreicht. Diese Transformation muss nicht explizit durchgeführt werden, sondern wird durch den sogenannten Kernel-Trick erreicht. Wie beschrieben, hat sich die RBF-Kernfunktion als besonders nützlich für diverse Anwendungsgebiete herausgestellt, darunter auch die Landnutzungsklassifikation. Im vorliegenden Abschnitt wurde ebenfalls beschrieben, wie alle Parameter einer SVM automatisch angelernt werden können, so dass ein optimales Klassifikationsergebnis erreicht werden kann. Vergleiche der SVM zu anderen Klassifikationsverfahren sowie die herausgearbeiteten Vorteile machen deutlich, dass die SVM als Klassifikationsalgorithmus besonders zum Trennen zweier Klassen mittels einer objektbasierten Klassifikation geeignet ist, auch dann, wenn räumlich getrennte Cluster der Klassen im Merkmalsraum vorliegen.

3.2. Einführung in die Segmentierung mit Wasserscheidentransformation und Regionennachbarschaftsgraph

Wie bereits herausgearbeitet, ist die Segmentierung von GIS-Objekten in unterschiedliche Segmente, korrespondierend zu unterschiedlichen Schlägen, ein Teilproblem der Verifikation von GIS-Acker- und Grünlandobjekten. Wie in Abschnitt 2.3 diskutiert, eignet sich zur Segmentierung ein Verfahren, basierend auf einer ursprünglich stark übersegmentierenden Wasserscheidentransformation mit anschließender Verschmelzung der Regionen mit ähnlichen Eigenschaften mittels eines RAG. Da es sich bei beiden Verfahren um weit verbreitete Standardverfahren der Bildverarbeitung handelt, werden die Grundlagen der Wasserscheidentransformation und des RAG in diesem Abschnitt nur kurz vorgestellt. Aufgeführte Erläuterungen wurden im wesentlichen Tönnies (2005) entnommen.

Die Wasserscheidentransformation basiert auf Gradienteninformationen und findet automatisch die Kanten eingeschlossener Segmente. Im topographischen Sinn trennt eine „Wasserscheide“ Gebiete, die in unterschiedliche Senken entwässert werden. Die Wasserscheiden werden im Bild so simuliert, dass sie

gerade dort verlaufen, wo zwei Bildsegmente voneinander unterschieden werden sollen. Dabei eignen sich Gradienten zum Simulieren der Topographie. Zur Rauschunterdrückung wird bei der Berechnung des Gradientenbildes eine Glättung, z.B. mittels Gaußfilter G_σ mit Glättungsparameter σ , durchgeführt. Die Wasserscheiden können anschließend mittels Beregnung oder Flutung bestimmt werden. Bei der zweiten Methode wird von den lokalen Minima aus geflutet, wie in Abbildung 5 zu sehen ist. Anschließend wird die Fluthöhe langsam gesteigert. Bei jedem Schritt wird für jedes neu überflutete Pixel überprüft, ob es

1. isoliert ist,
2. eine Überflutungsregion steigert oder
3. eine Wasserscheide ist.

Ist ein neu überflutetes Pixel isoliert, also ist es nicht zu anderen überfluteten Pixeln benachbart, dann handelt es sich um ein neu gefundenes Minimum. In diesem Fall wird für dieses isolierte Pixel ein neues Segment im Segmentierungsbild erzeugt. Steigert ein neu überflutetes Pixel eine Überflutungsregion, wird dieses Pixel dem entsprechenden Segment zugeordnet. Ist das neu überflutete Pixel eine Wasserscheide, d.h. ist dieses Pixel von mindestens zwei Regionen benachbart, ist die Segmentgrenze erreicht.

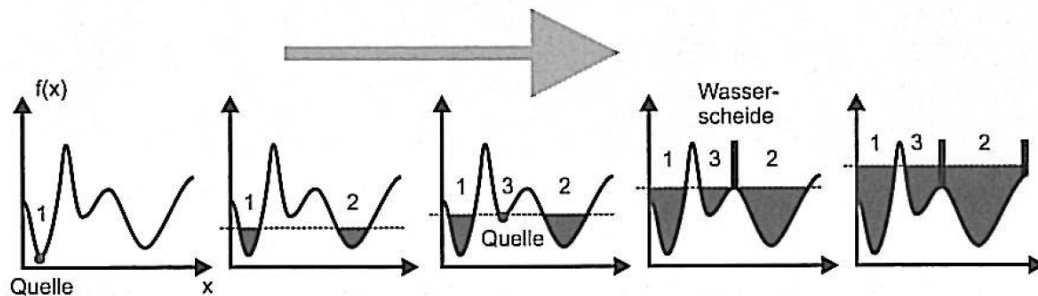


Abbildung 5: Schema der Wasserscheidentransformation durch Flutung, Abbildung übernommen aus (Tönnies, 2005), S. 225.

Der Grad der Übersegmentierung, der bei der Wasserscheidentransformation erreicht wird, hängt vom Grad der Glättung des Gradientenbildes ab. Detaillierte Informationen und Beschreibungen zur Wasserscheidentransformation können neben Tönnies (2005) auch Gonzalez und Woods (2002) entnommen werden.

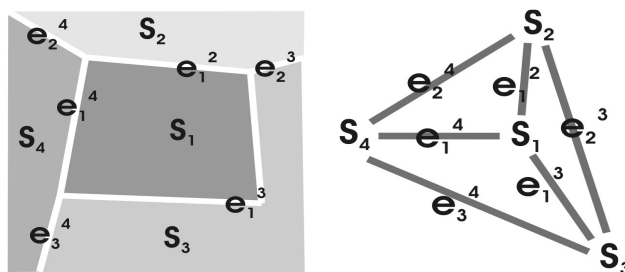


Abbildung 6: RAG (rechts) eines synthetischen Bilds (links)
(S_i ... Segmente, e_i^j ... Kanten zwischen den Segmenten S_i und S_j).

Nach der Wasserscheidensegmentierung wird ein RAG generiert. Die Knoten des Graphen sind die Segmente, während die Kanten die Nachbarschaftsrelationen repräsentieren: zwei Segmente S_i und S_j mit $i \neq j$ sind durch die Kante e_i^j im RAG verbunden, wenn es nach der Segmentierung mindestens einen Grenzpixel gibt, das sowohl an S_i als auch an S_j grenzt. Ein Beispiel eines solchen Graphen für ein synthetisches Bild ist Abbildung 6 zu entnehmen.

In der Regel werden, wenn der RAG konstruiert wird, gleichzeitig Attribute für Knoten und die Kanten berechnet. Attribute können z.B. geometrische oder radiometrische Attribute sein. Der RAG mit den Attributen für Knoten und Kanten ist Voraussetzung für das Verschmelzen benachbarter Regionen zur Verbesserung der initialen Segmentierung. Ziel des Prozesses ist es, Regionen zu verschmelzen, die gemeinsame Eigenschaften besitzen. Diese Eigenschaften können unterschiedlich definiert werden, genauso wie die Kriterien, unter denen zwei Segmente nicht verschmolzen werden dürfen. Eigenschaften und Kriterien sind entsprechend der Anwendung zu wählen.

4. Neue Methode zur Verifikation der Objektarten Ackerland und Grünland eines GIS

In diesem Kapitel wird ein neues Verfahren zur Verifikation der Objektarten Acker- und Grünland eines GIS eingeführt. Nachdem die Eingangsdaten erläutert wurden, wird zunächst ein Überblick über den Algorithmus gegeben, bevor auf die einzelnen Schritte für die Verifikation der Objektarten Ackerland und Grünland eingegangen wird. Das Kapitel schließt mit einer Bewertung des Algorithmus inklusive seiner Grenzen.

4.1. Eingangsdaten

Als Eingangsdaten in das Verfahren zur Verifikation von GIS-Objekten der Klassen Acker- und Grünland dienen:

- das zu verifizierende GIS und
- ein Bilddatensatz.

Der in dieser Arbeit vorgestellte Ansatz ist auf Spezifikationen von GIS-Datenbanken ausgerichtet, die den Inhalt und Detaillierungsgrad einer topographischen Karte mittleren Maßstabs entsprechen, wie z.B. das ATKIS Basis-DLM (AdV, 1997), CLC (EEA, 2011) und MGCP.

Der Bilddatensatz besteht aus einem monotemporalen Luft- oder Satellitenbild mit einer geometrischen Auflösung von 0,5 bis 1 m. Das Bild muss orthorektifiziert und multispektral (in Fall von Satellitenbildern pansharpend) sein, wobei die Farbinformation der Kanäle *Rot* (R), *Grün* (G) und *Blau* (B) und/oder der Kanäle *Nahes Infrarot* (NIR), *Rot* (R) und *Grün* (G) vorliegen können.

4.2. Überblick

Ziel ist die automatisierte Verifikation von GIS-Objekten der Klassen Acker- und Grünland. Die automatisierte Verifikation soll den menschlichen Bearbeiter bei der Qualitätskontrolle unterstützen. Dabei sollen zum einen genügend Fehler gefunden werden, so dass die geforderten Qualitätsansprüche erfüllt werden, zum anderen soll dem Bearbeiter bei Benutzung des Systems eine Zeitersparnis entstehen. Wie in Kapitel 1 herausgearbeitet, wird in dieser Arbeit ein semi-automatisches Verifikationsverfahren verwendet. Alle GIS-Objekte, die vom System als richtig verifiziert werden, brauchen vom menschlichen Bearbeiter nicht mehr betrachtet zu werden. Alle anderen Objekte werden manuell überprüft. Die endgültige Entscheidung über Annahme oder Rückweisung eines vom System abgelehnten Objektes wird von einem menschlichen Bearbeiter durchgeführt. Dieser Schritt wird in der Nachbearbeitung des Verfahrens ausgeführt.

Wie in der Diskussion in Abschnitt 2.4 herausgearbeitet, wird für die Verifikation von Ackerland- und Grünlandobjekten ein zweistufiges Verfahren verwendet. Im ersten Schritt wird die gemeinsame Klasse *Acker-/Grünland* gegenüber allen anderen Klassen wie z.B. *Siedlung*, *Industrie* und *Wald* mit Hilfe einer pixelbasierten Klassifikation abgegrenzt, während im zweiten Schritt *Acker-* und *Grünland* mit Hilfe einer objektbasierten Klassifikation voneinander getrennt werden. Durch die zweistufige Vorgehensweise ist es möglich, die Vorteile einer pixelbasierten und einer objektbasierten Klassifikation zu verbinden, wie im Abschnitt 2.1.1.3. diskutiert. Die Literaturübersicht (Kapitel 2) zeigte, dass ein zweistufiges Verfahren, das eine pixel- und objektorientierte Klassifikation verbindet, bisher weder für die Verifikation noch für die Klassifikation landwirtschaftlicher Flächen zum Einsatz kam.

In Kapitel 2 zeigte sich, dass für den ersten Schritt, der Trennung der gemeinsamen Klasse *Acker-/Grünland* von weiteren Klassen, ein pixelbasiertes Verfahren am besten geeignet ist, da es auf diese Weise möglich ist, auch einzelne kleine Flächen mit Fremdnutzung, die in Ackerland- und Grünlandobjekten nicht erlaubt sind, zu detektieren. In der Diskussion in Abschnitt 2.4 zeigte sich, dass für eine pixelbasierte Klassifikation von mehreren Klassen das MRF-Verfahren besonders geeignet ist. So kann das Verfahren der MRF zum einen Kontext berücksichtigen, was bei einer pixelbasierten Klassifikation zur Reduktion des „Salt-and-Pepper-Effects“ führt; und zum anderen erlaubt das MRF-Verfahren die Klassifikation einer beliebigen Anzahl von Klassen. Ein solcher Ansatz wurde in (Busch et al., 2005) benutzt. Busch et al. (2005) zeigen, dass mittels des MRF-Ansatzes eine zuverlässige Trennung der gemeinsamen Klasse *Acker-/Grünland* gegenüber anderen Klassen möglich ist. Es wurde auch gezeigt, dass eine Trennung von *Acker-* und *Grünland* mit diesem Ansatz nicht durchgeführt werden kann. Die Arbeiten von Busch et al. (2005) basieren auf einem Ansatz von

Gimel'farb (1996). Nach der pixelbasierten Klassifikation werden deren Ergebnisse auf das GIS-Objekt übertragen, wobei die Spezifikationen des GIS berücksichtigt werden. Ergibt sich daraus, dass das Objekt weder der Objektart Ackerland noch der Objektart Grünland angehört, wird das Objekt vom System als falsch gekennzeichnet, andernfalls wird es zum zweiten Analyseschritt weitergeleitet. In der vorliegenden Arbeit wird das Verfahren nach Busch et al. (2005), das auf den Arbeiten von Gimel'farb (1996) beruht, für die Trennung der gemeinsamen Klasse *Acker-/Grünland* von anderen Klassen wie *Siedlung*, *Industrie* und *Wald* angewendet, da es für diese Aufgabe bereits erfolgreich eingesetzt wurde. Da auf bereits existierende Verfahren zurückgegriffen wird, werden in diesem Kapitel diese Verfahren nur kurz vorgestellt. Nähere Informationen sind Gimel'farb (1996) und Busch et al. (2005) zu entnehmen.

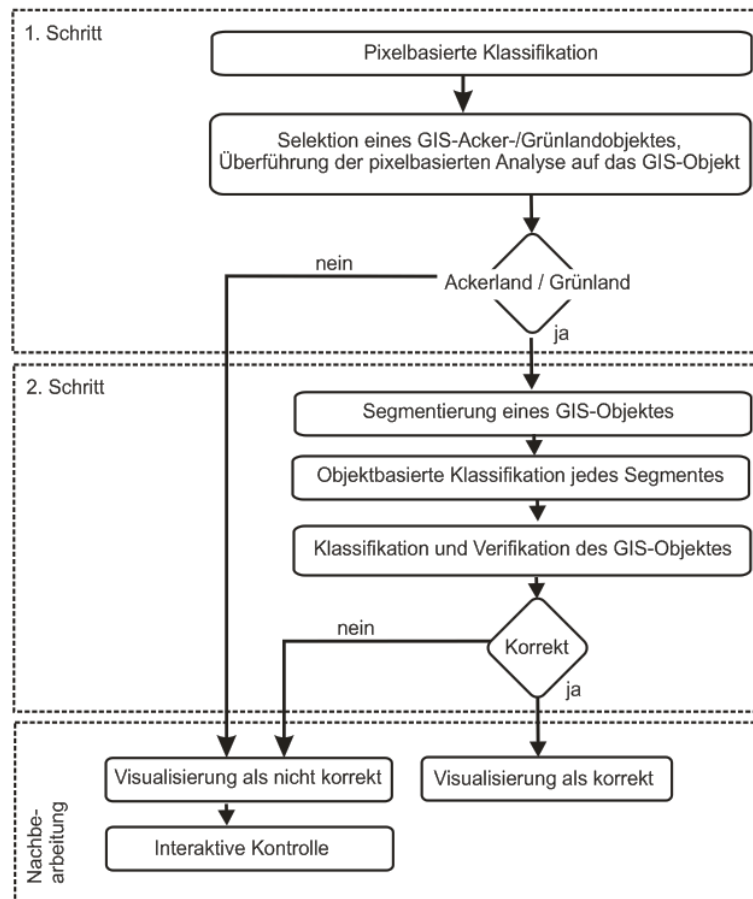


Abbildung 7: Ablaufdiagramm der Methode zur Verifikation der Objektarten Ackerland und Grünland eines GIS.

In einem zweiten Schritt werden alle GIS-Objekte, die im ersten Schritt nicht abgelehnt worden, in die Objektarten Acker- und Grünland klassifiziert. Das Verfahren des zweiten Analyseschrittes ist neuartig und wird daher in diesem Kapitel im Detail beschrieben.

Kapitel 2 und Kapitel 3 zeigten, dass sich für den zweiten Analyseschritt das SVM-Klassifikationsverfahren besonders gut eignet, da die SVM zum einen mit räumlich getrennten Clustern von Klassen im Merkmalsraum umgehen kann, wie sie für die Klasse Ackerland auf Grund der unterschiedlichen Erscheinungsformen zu erwarten sind, und zum anderen ein binäres Klassifikationsverfahren ist. Wie Kapitel 2 zeigte, wurde das Verfahren der SVM zur Klassifikation landwirtschaftlicher Flächen bisher nicht genutzt. Die Grundlagen wurde in Abschnitt 3.1 detailliert eingeführt.

Bevor die Klassifikation innerhalb des zweiten Analyseschrittes vorgenommen werden kann, muss eine Segmentierung der einzelnen Schläge innerhalb eines GIS-Objektes erfolgen. In Kapitel 2 zeigte sich, dass dafür die Wasserscheidentransformation mit anschließender Verwendung eines RAG geeignet ist. Dabei werden nach einer ursprünglich stark übersegmentierenden Wasserscheidentransformation benachbarte Regionen mit ähnlichen Eigenschaften verschmolzen. Aus diesem Grund wird dieses Verfahren in der vorliegenden Arbeit genutzt; in diesem Kapitel werden die Kriterien vorgestellt, nach denen die Regionen des

RAG verschmolzen werden. Diese Kriterien sind speziell für die in dieser Arbeit geforderten Kriterien zur Segmentierung von Schlägen ausgerichtet. Nach der Segmentierung erfolgt zunächst eine Zuordnung jedes Segmentes eines GIS-Objektes zu der Objektart Ackerland oder Grünland. Sind alle Segmente eines GIS-Objektes einer Klasse zugewiesen, erfolgt die Klassifikation und Verifikation des GIS-Objektes. Auf Grund der Verwendung einer automatischen Segmentierung zur Bestimmung einzelner Schläge, deren getrennter Klassifikation und gemeinsamer Verifikation, hebt sich die vorliegende Arbeit von bisher existierenden Klassifikationsverfahren ab (Kapitel 2). Diese Lösung setzt somit kein Schlagkataster voraus, wie es existierende Verfahren zur Klassifikation von Acker- und Grünland i.d.R. benötigen.

Für die objektbasierte Klassifikation von Ackerland- und Grünlandsegmenten sind Merkmale spektrale, texturelle, strukturelle und geometrische Merkmale nötig, wie in Abschnitt 2.4 diskutiert. Dort wurde ebenfalls festgestellt, dass bei den bisher verwendeten strukturellen Merkmalen häufig Vorinformationen benötigt werden, die bei Objektklassen wie *Weingarten* einfacher verfügbar und meist bekannt sind, jedoch nicht bei Ackerlandflächen. Für die Trennung von Ackerland gegenüber Grünland wird daher in diesem Kapitel ein Ansatz zur Bestimmung struktureller Merkmale vorgestellt, der keine Vorinformation benötigt.

Werden alle bisher in diesem Abschnitt genannten Punkte berücksichtigt, ergibt sich ein Ablauf der Methode wie in Abbildung 7 zusammengefasst. Im weiteren Verlauf dieses Kapitels werden die einzelnen Schritte, die im Ablaufdiagramm in Abbildung 7 zu sehen sind, beschrieben.

Zusammenfassend kann festgestellt werden: Die neue Methode zur Verifikation von Ackerland- und Grünlandobjekten wird in der Lage sein, die Verifikation dieser GIS-Objekte auf Grundlage einer Klassifikation der getrennten Klassen Acker- und Grünland durchzuführen. Damit wird erstmals die Möglichkeit der automatisierten Verifikation von GIS-Ackerland- und Grünlandobjekten bestehen. Dieses Verfahren für die Verifikation berücksichtigt dabei in Klassifikation und Verifikation das Vorhandensein mehrerer Schläge innerhalb eines GIS-Ackerland- und Grünlandobjektes mittels automatischer Segmentierung. Das findet bei bisher existierenden Klassifikationsverfahren keine Berücksichtigung. Bei der Klassifikation der Segmente werden z.T. extra für die Problematik der Trennung von Acker- und Grünland entwickelte Merkmale verwendet, die neuartig sind und detailliert vorgestellt werden. Außerdem werden Teilnutzungen in Ackerland- und Grünlandobjekten berücksichtigt, was bisher bei Klassifikationsansätzen, wie im Kapitel 2 gezeigt, ebenfalls keine Berücksichtigung fand.

4.3. Klassifikation zur Identifizierung von GIS-Ackerland- und Grünlandobjekten

Im ersten Schritt des zweistufigen Prozesses zur Verifikation von Ackerland- und Grünlandobjekten eines GIS werden zunächst *Acker-* und *Grünland* zusammen als eine Klasse betrachtet und gegen alle anderen Klassen abgegrenzt. Andere Klassen können *Siedlung*, *Industrie* und *Wald* sein. Diese Klassifikation soll auch in der Lage sein, einzelne Objekte, die einer anderen Objektart angehören, innerhalb eines GIS-Objektes zu detektieren, wie z.B. Häuser, die in GIS-Ackerland- oder Grünlandobjekten nicht zulässig sind. Ein Beispiel für diesen Fall ist in Abbildung 8 zu sehen. Hier entstanden innerhalb von vier Jahren mehrere Häuser eines neuerschlossenen Wohngebietes (durch Pfeile gekennzeichnet), wo zuvor ein reines GIS-Grünlandobjekt existierte. Es ist zu sehen, dass das GIS zum späteren Zeitpunkt noch nicht aktualisiert ist und daher einen Fehler beinhaltet (Häuser in einem Grünlandobjekt).

Busch et al. (2005) haben nachgewiesen, dass ihre Methode für die notwendigen Anforderungen geeignet ist, weshalb diese Methode auch in dieser Arbeit verwendet wird. Die Methode von Busch et al. (2005) wurde ebenfalls in den Arbeiten von Müller (2007) erfolgreich eingesetzt, bei der ähnliche Anforderungen, bezogen auf die Objektklassen *Siedlung* und *Industrie* bestanden. Die in der Arbeit von Busch et al. (2005) verwendete pixelbasierte Klassifikation basiert auf dem Ansatz von Gimel'farb (1996), während die Überführung des pixelbasierten Klassifikationsergebnisses auf GIS-Objekte von Busch et al. (2005) selber stammt. In diesem Absatz werden beide Methoden (pixelbasierte Klassifikation nach Gimel'farb (1996) und Überführung des pixelbasierten Klassifikationsergebnisses nach Busch et al. (2005)) nur kurz vorgestellt, da es sich um bereits existierende Verfahren handelt. Nähere Informationen sind Gimel'farb (1996, 1997) und Busch et al. (2004, 2005) zu entnehmen.

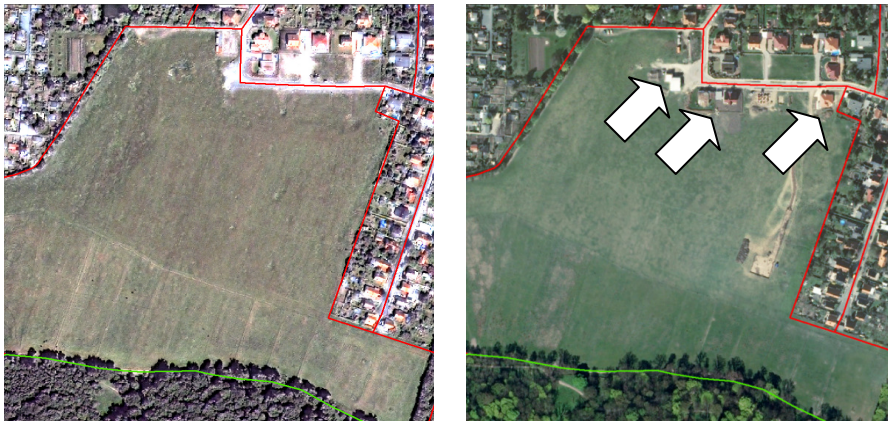


Abbildung 8: Grünlandobjekt in 2006 ohne Häuser (links, IKONOS) und in 2010 mit mehreren Häusern (rechts, Luftbild, mit Pfeilen gekennzeichnet), die Ausschnitte zeigen ein jeweils 500 m x 500 m großes Gebiet, farbige Linien entsprechen den Grenzen von GIS-Objekten.

4.3.1. Pixelbasierte Klassifikation

Das Verfahren von Gimel'farb (1996) beruht auf der Analyse von Textur und Farbe. Es handelt sich um eine überwachte Klassifikation, basierend auf MRF. Wie bereits im Abschnitt 2.1.1.2 angemerkt, wird das Bild bei MRF als ungerichteter Graph interpretiert, bei dem die Knoten den Pixeln und die Kanten den Beziehungen zwischen benachbarten Pixeln entsprechen. Auf Grund der Modellierung von Interaktionen benachbarter Pixel bei der Klassifikation können lokale Zusammenhänge berücksichtigt werden. Das zu klassifizierende Pixel wird damit nicht nur auf Grund seiner Merkmale, sondern auch auf Grund der Klassenzugehörigkeit aller anderen Pixel über eine Wahrscheinlichkeitsfunktion einer Klasse zugeordnet.

MRF sind, wie in Kapitel 2 gezeigt, gut für die pixelbasierte Klassifikation geeignet, da es auf Grund der Verwendung des lokalen Kontextmodells einer Glättung des Klassifikationsergebnisses kommt und somit der „Salt-and-Pepper-Effect“ bei den unabhängigen Klassen von Pixeln verringert wird. Ein Vorteil des MRF-Verfahrens gegenüber der SVM ist, dass das MRF-Verfahren grundsätzlich für die Klassifikation mehrerer Klassen geeignet ist, wie es sich hier als notwendig erweist. Die zu bestimmenden Objektklassen sind unter Berücksichtigung der in der Szene charakteristischen Klassen zu wählen. Charakteristische Klassen für ländliche Regionen sind *Acker-/Grünland, Siedlung, Industrie* und *Wald*.

Für die Klassifikation in (Gimel'farb, 1996) werden Texturmerkmale verwendet, die durch die Interaktion von Pixelpaaren modelliert werden. Für die Modellierung der Interaktionen werden Gibbs-Potenziale verwendet. Die Interaktion von Pixelpaaren drückt dabei aus, dass die Grauwerte innerhalb einer lokalen Nachbarschaft nicht zufällig verteilt sind, sondern einer bestimmten Konfiguration unterliegen. Die für die Klassifikation verwendeten Merkmale beruhen auf Grauwertdifferenzhistogrammen (*Gray Level Difference Histograms* – GLDH), die in einer 33 x 33 Pixel großen Umgebung für den zentralen zu klassifizierende Pixel berechnet werden. Die GLDH wird nicht nur pro Bildkanal, sondern auch zwischen jeweils zwei Kanälen bestimmt. Da die lokale Umgebung zur Berechnung der GLDH auf 33 x 33 Pixel festgelegt ist, wird, um die Skalenabhängigkeit der Textur zu berücksichtigen, eine Anpassung der geometrischen Auflösung der Eingangsbilder an die zu untersuchenden Texturen vorgenommen. Dafür wird durch Unterabtastung der Eingabebilddaten eine Bildpyramide mit L Auflösungsstufen erzeugt. Um Aliasing-Artefakte zu vermeiden, wird das Eingabebild vor Unterabtastung tiefpassgefiltert. In jedem Schritt der Unterabtastung wird die geometrische Auflösung des Bildes um den Faktor 2 reduziert. Die Anzahl der Auflösungsstufen in dieser Arbeit wurde auf 4 festgelegt. Mittels Training wird für jede Klasse die geeignete Auflösungsstufe bestimmt und für die Klassifikation angewendet. Dafür wird das Verfahren von Becker et al. (2009) eingesetzt. Um die Auflösungsstufe zu bestimmen, trainiert Becker et al. (2009) zunächst das MRF pro Auflösungsstufe. Anschließend findet pro Auflösungsstufe die Reklassifizierung der Trainingsdaten statt. Durch Vergleich der Reklassifizierungsergebnisse jeder Auflösungsstufe erhält man schließlich die beste Auflösungsstufe je Klasse. Durch die Lernstichprobe wird das Training manuell festgelegt.

4.3.2. Überführung der pixelbasierten Klassifikation auf ein GIS-Objekt

Da es sich bei dem in (Gimel'farb, 1996) vorgestellten Ansatz um ein pixelbasiertes Verfahren handelt, die Verifikation aber pro GIS-Objekt durchzuführen ist, gilt es, das pixelbasierte Verfahren auf das GIS-Objekt zu übertragen und dabei Fehler, wie z.B. neu entstandene Häuser (Abbildung 8), zu detektieren. Das Verfahren zur Übertragung des pixelbasierten Klassifikationsergebnisses auf ein GIS-Objekt wurde (Busch et al., 2005) entnommen.

Pixel, die mit der Objektklasse im GIS konsistent sind, werden als „korrekt“, alle anderen Pixel als „falsch“ gekennzeichnet. Insgesamt werden zwei Kriterien für die Bewertung eines GIS-Objektes herangezogen. Das erste Kriterium ist der Quotient q aus der Anzahl von falschen Pixeln und der Anzahl aller Pixel, die das GIS-Objekt bedecken.

$$q = \frac{N_{falsch}}{N_{korrekt} + N_{falsch}} \quad (4.1)$$

Auf Grund von Rauschen und inhomogenen Strukturen können falsche Pixel mehr oder weniger gleichmäßig über ein Objekt verteilt sein. Es ist z.B. möglich, dass bei der Klassifikation einer Siedlungsfläche in offener Bauweise große Dachflächen als *Industrie* sowie Bäume und Wiesen in Hintergärten als *Wald* oder *Acker-/Grünland* detektiert werden, wie in Abbildung 9 sichtbar. Wenn es sich bei diesen Flächen jedoch um kompakte Fehler handelt und nicht um Rauschen, muss das GIS-Objekt vom System als falsch abgelehnt werden, auch wenn der Schwellwert für Quotienten q (4.1) nicht erreicht ist. Ein kompakter Fehler ist:

- ein Verbund von falschen Pixeln,
- die ein und derselben Klasse angehören sowie
- eine gewisse Größe und Breite besitzen.

Das Maß der Breite, das hier Anwendung findet, ist die Anzahl der Schritte eines morphologischen Filters (Erosion) E , die benötigt werden, bis die Region des vermeintlichen kompakten Fehlers verschwindet. Ist die Anzahl höher als ein definierter Schwellwert s_E und besitzt der Fehler eine Mindestfläche s_A , handelt es sich um einen kompakten Fehler. Liegt mindestens ein kompakter Fehler vor, wird das GIS-Objekt vom System abgelehnt.



Abbildung 9: Klassifikationsergebnis (rechts) eines Siedlungsgebietes (links) - Siedlung (rot), Industrie (gelb), Wald (dunkelgrün), Acker-/Grünland (hellgrün).

Im Gegensatz zu q , wo die Summe über alle falsch klassifizierte Pixel gebildet wird, egal welcher Klasse diese angehören, betrachtet der kompakte Fehler nur Pixel, die zu ein und derselben Klasse gehören. Der für q zu wählende Schwellwert s_q sowie die Größe des Schwellwertes s_A für die Detektion eines kompakten Fehlers ergeben sich aus den Spezifikationen des jeweiligen GIS-Objektartenkataloges. Der Schwellwert s_E für die Bestimmung eines kompakten Fehlers ergibt sich z.B. aus dem Vorwissen über die Dimension eines Hauses (Abbildung 8).

Wenn ein oder mehrere Gebäude in einem GIS-Grünlandobjekt neu entstanden sind, kann das Neubaugebiet das komplette oder einen sehr großen Teil des Ackerlandobjektes bedecken. In diesem Fall wird der Fehler im GIS sowohl über q als auch über das Vorhandensein eines kompakten Fehlers der Klasse Siedlung entdeckt. Bedeckt das Neubaugebiet nur einen Teil des Ackerlandobjektes, dann wird dieser Fehler als kompakter Fehler erkannt. Wenn kein kompakter Fehler gefunden werden konnte, weil es sich z.B. um ein

sehr langes und schmales GIS-Objekt bzw. um ein sehr langes und schmales Neubaugebiet handelt, kann der Fehler im GIS fast immer über q gefunden werden.

Zusammenfassend kann gesagt werden, dass, wenn nach der Klassifikation des Ansatzes nach Gimel'farb (1996) ein GIS- Ackerland-/Grünlandobjekt

- aus Pixeln besteht, die der Klasse Acker-/Grünland angehören,
- q des GIS-Objektes kleiner als ein zuvor definierter Grenzwert s_q und
- die Anzahl der kompakten Fehler Null ist (also die Schwellwerte s_A und s_E nicht überschritten werden),

dann handelt es sich bei diesem GIS-Objekt tatsächlich um ein Ackerland- oder Grünlandobjekt, das zur nächsten Analysestufe weitergegeben wird. Ist eine dieser Bedingungen nicht erfüllt, ist ein Fehler im GIS nicht auszuschließen. Das GIS-Objekt wird in diesem Fall vom System abgelehnt und zur endgültigen Bewertung einem menschlichen Bearbeiter vorgelegt.

4.4. Segmentierung von Schlägen

Die Segmentierung von GIS-Objekten in Schläge ist notwendig, weil viele GIS, z.B. ATKIS, das Vorhandensein mehrerer Schläge innerhalb eines GIS-Objektes erlauben (Abbildung 10) und die vorgestellte Methode mit dieser Spezifikation umgehen können muss. Die verschiedenen Schläge erscheinen nach der Glättung mittels Gaußfilter G_σ mit dem Glättungsparameter σ als einigermaßen homogene Segmente in den Bildern. Die Schläge können also nach einer Gaußfilterung mittels Segmentierung als homogene Segmente extrahiert werden. Dafür wird ein Segmentierungsalgorithmus, wie in Abschnitt 3.2 eingeführt, verwendet.



Abbildung 10: Falschfarbeninfrarotbilder zweier GIS-Ackerlandobjekte, die aus mehreren Schlägen bestehen. GIS-Grenzen sind in cyan und blau eingezeichnet.

Zunächst wird eine Wasserscheidensegmentierung angewendet. Abbildung 11 zeigt ein RGB-Bild mit den Ergebnissen zweier Wasserscheidensegmentierungen unter Verwendung verschiedener Glättungsparameter. Die Wasserscheidensegmentierung liefert Regionen mit geschlossenen Grenzen, auch wenn das Bild in der Mitte eine starke Übersegmentierung zeigt. Die Segmentgrenzen, die ein menschlicher Bearbeiter wählen würde, sind alle vorhanden, aber es sind viele zusätzliche Grenzen enthalten. Im Segmentierungsbild auf der rechten Seite sind auf Grund des stärker geglätteten Eingabebildes wichtige Bildstrukturen verschmolzen. Eine gute Trennung der Schläge ist somit nicht mehr möglich. Um diesen Problemen zu begegnen, wird in dieser Arbeit eine iterative Lösung durchgeführt. Zunächst wird eine Wasserscheidensegmentierung mit einem geringen Glättungsgrad angewendet, das zu einer starken Übersegmentierung führt. Anschließend wird ein Regionennachbarschaftsgraph (*region adjacency graph*, RAG) generiert, der wichtige Attribute der Bildregionen und deren Grenzen enthält. Diese Attribute sind speziell für die geforderte Aufgabe entwickelt worden und werden daher in diesem Abschnitt detailliert eingeführt. Nachdem die Regionen, basierend auf den Attributen, verschmolzen wurden, wird abschließend eine Nachbearbeitung der Segmentierungsergebnisse an Hand von GIS-Spezifikationen durchgeführt.

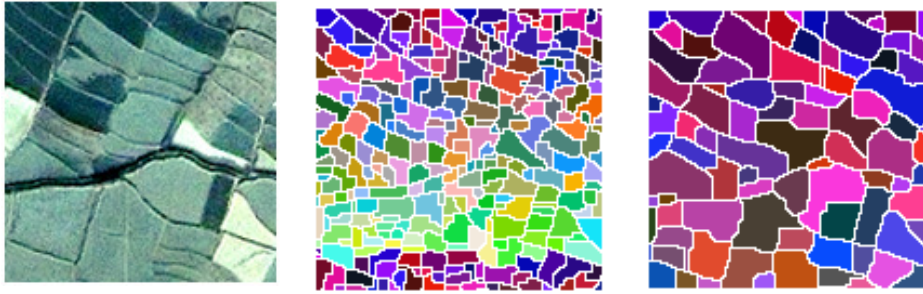


Abbildung 11: Wasserscheidensegmentierung eines RGB-Bildes (links) für $\sigma = 3$ (Mitte) und $\sigma = 8$ (rechts).

4.4.1. Attribute der Bildregionen und deren Grenzen

Wenn der RAG konstruiert wird, werden gleichzeitig die Attribute für Knoten und die Kanten berechnet. Eine Region hat geometrische Attribute und radiometrische Attribute.

Die geometrischen Attribute eines Segmentes i sind:

- Anzahl der Pixel N_i ,
- Minimum- und Maximumkoordinaten,
- Koordinaten des Schwerpunktes.

Zur Berechnung der radiometrischen Attribute nimmt man zunächst an, dass ein multispektrales Bild I mit N_k Kanälen durch die Grauwertvektoren $\mathbf{g}(x, y) = [g_1(x, y), g_2(x, y), \dots, g_{N_k}(x, y)]^T$ an der Stelle (x, y) repräsentiert wird. Die radiometrischen Attribute eines Segmentes i sind:

- mittlerer Grauwertvektor $\mathbf{g}_{avg}^i$,
- Kovarianzmatrix $\mathbf{Q}_{gg}^i$ der Grauwerte,
- Gesamtmaß var^i des Rauschniveaus innerhalb des Segments als Spur von $\mathbf{Q}_{gg}^i$ mit $var^i = Spur(\mathbf{Q}_{gg}^i)$.

Um die Berechnung von $\mathbf{g}_{avg}^i$ und $\mathbf{Q}_{gg}^i$ robust in Hinsicht auf Ausreißer nahe der Segmentgrenze zu machen, werden Grauwertvektoren nahe der Segmentgrenze nicht berücksichtigt. Es werden nur dann alle Grauwertvektoren verwendet, wenn ein Segment so klein ist, dass alle Pixel nahe einer Grenze liegen und daher alle Pixel aus der Berechnung ausgeschlossen werden müssten.

Eine Kante e_i^j im RAG repräsentiert die Nachbarschaftsbeziehung zwischen zwei Segmenten S_i und S_j und die Grenze zwischen zwei Regionen. Eine Grenze zwischen zwei Regionen besteht aus einem oder einer Sequenz von Grenzpixeln. Daher besteht jede Kante im RAG aus einem Satz von verbundenen Grenzpixeln, die aus dem Labelbild, das das Segmentierungsergebnis repräsentiert, extrahiert werden. Weiterhin besitzen die Grenzpixel eine 2D Ausdehnung im digitalen Bild, nämlich die Fläche, die von diesen Pixeln bedeckt wird. Aus diesem Grund besitzt die Kante e_i^j ebenfalls einen durchschnittlichen Grauwertvektor und eine Kovarianzmatrix der Grauwerte, berechnet aus den Grauwerten aller Grenzpixel, die S_i und S_j trennen.

4.4.2. Verschmelzen von Regionen

Der RAG mit den Attributen für Knoten und Kanten ist Voraussetzung für das Verschmelzen benachbarter Regionen zur Verbesserung der initialen Segmentierung. Ziel des Prozesses ist es, Regionen zu verschmelzen, die

- ähnliche radiometrische Eigenschaften besitzen und
- durch keine signifikante Kante getrennt sind.

Ähnliche radiometrische Eigenschaften können mit Hilfe eines Hypothesentest erkannt werden. Dafür wird die Differenz der beiden mittleren Grauwertvektoren zweier Regionen Δ_{ij} mit $\Delta_{ij} = (\mathbf{g}_{avg}^i - \mathbf{g}_{avg}^j)$ mit $\mathbf{Q}_{\Delta\Delta}^{ij} = \frac{1}{N_i} \mathbf{Q}_{gg}^i + \frac{1}{N_j} \mathbf{Q}_{gg}^j$ mittels statistischen Tests daraufhin untersucht, ob Δ_{ij} signifikant von Null verschieden ist. Dies führt auf die Testgröße

$$D_{ij} = \Delta_{ij}^T \mathbf{Q}_{\Delta\Delta}^{ij-1} \Delta_{ij} \sim \chi_{N_k}^2, \quad (4.2)$$

wobei D_{ij} unter der Annahme der Normalverteilung der Grauwerte der Chi-Quadrat-Verteilung unterliegt und N_k der Anzahl der Freiheitsgrade (hier: Anzahl der Kanäle) entspricht. Experimente haben gezeigt, dass D_{ij} in Gleichung (4.2) ein zu strenges Kriterium ist. Daher wird

$$D_{ij} = \Delta_{ij}^T (\mathbf{Q}_{gg}^i + \mathbf{Q}_{gg}^j)^{-1} \Delta_{ij} \quad (4.3)$$

verwendet. Dieses benutzt die Differenzen der ursprünglichen Grauwertvektoren und nicht die Mittelwerte für den Hypothesentest.

Beim Test eines zweiten Maßes werden zwei Modelle verglichen: Es wird untersucht, ob die Gesamtvarianz der Grauwerte der verschmolzenen Regionen signifikant größer ist als die Varianz beider separater Regionen. Zum ersten wird die Gesamtvarianz für das Modell zweier getrennter Regionen var_s^{ij} mit unterschiedlichen Grauwertmittelwerten mit

$$var_s^{ij} = \frac{(N_i \cdot var^i) + (N_j \cdot var^j)}{(N_i + N_j - 1)} \quad (4.4)$$

bestimmt, wobei N_i und N_j der Anzahl der Pixel der Regionen S_i und S_j entsprechen. Zum anderem wird die Varianz der verschmolzenen Regionen var_m^{ij} mit

$$var_m^{ij} = \text{Spur} (\mathbf{Q}_{gg}^m) \quad (4.5)$$

berechnet. Beide Varianzen werden nun verglichen:

$$F_{ij}^v = \frac{var_m^{ij}}{var_s^{ij}} \sim F_{N_k(N_i-1)+N_k(N_j-1), N_k(N_i+N_j-1)}, \quad (4.6)$$

wobei F_{ij}^v unter der Annahme von Normalverteilung der Grauwerte nach Fischer verteilt ist und $N_k(N_i - 1) + N_k(N_j - 1)$ sowie $N_k(N_i + N_j - 1)$ die Anzahl der Freiheitsgrade entspricht. Die Regionen können nur dann verschmolzen werden, wenn F_{ij}^v kleiner als ein definierter Schwellwert s_F ist, also wenn var_m^{ij} für das Modell der verschmolzenen Regionen nicht signifikant größer als var_s^{ij} unter der Annahme des Modells der separaten Regionen ist.

Schließlich sollen zwei Regionen auch dann nicht verschmolzen werden, wenn zwischen ihnen eine signifikante Linie besteht, auch wenn die Regionen eine ähnliche Grauwertverteilung besitzen. Eine signifikante Linie kann z.B. ein Weg sein. Die Stärke einer Linie wird über das Maß T_{ij} eingeführt. T_{ij} gibt den Anteil von Grenzpixeln an, die einen Homogenitätswert H (Förstner, 1994) größer als einen Schwellwert H_{max} besitzen. H basiert auf der ersten Ableitung der Grauwerte in der lokalen Nachbarschaft und berechnet sich aus:

$$H = \sum_{i=1}^{N_k} \frac{G_\sigma * (\Delta g_{ix}^2 + \Delta g_{iy}^2)}{\sigma_{ni}^2}, \quad (4.7)$$

wobei Δg_{ix} und Δg_{iy} die ersten Ableitungen der Grauwerte g_i von Band i jeweils nach x und y sind. G_σ ist ein Gaußfilter mit dem Glättungsparameter σ , wobei σ_{ni}^2 die Varianz der geglätteten Grauwertdifferenzen $G_\sigma * \Delta g_{ix}$ und $G_\sigma * \Delta g_{iy}$ ist, welches von der Schätzung der Rauschvarianz σ_{ni}^2 von Kanal i abgeleitet wird (Brügelmann und Förstner, 1992). Die Summe wird über alle N_k Kanäle des Bildes gebildet.

T_{ij} kann als Prozentzahl der Kantenpixel der Grenze, die zwei Regionen voneinander trennen, interpretiert werden. Er wird groß sein, wenn die Grenze zu einer signifikanten Linie im Bild korrespondiert und demnach zu einer wirklichen Grauwertdiskontinuität gehört.

Zwei Regionen sollen nur verschmolzen werden, wenn

- D_{ij} aus Gleichung (4.3) kleiner als ein Schwellwert s_D ,
- F_{ij} aus Gleichung (4.6) kleiner als ein Schwellwert s_F und
- T_{ij} kleiner als ein Schwellwert s_T

ist. Gleiche Bedingungen gelten auch für die Kanten, da wie oben festgelegt, eine Kante des RAG auch einen Vektor mit Grauwerten und so auch eine Kovarianzmatrix der Grauwerte besitzt.

Die vom Nutzer festzulegenden Schwellwerte s_D und s_F sind nicht leicht interpretierbar. Eine leichtere Interpretierbarkeit der einzustellenden Parameter ist jedoch wünschenswert. Aus diesem Grund formulieren wir s_D und s_F mit

$$s_D = k_D \cdot \chi_{N_k, 1-\alpha} \quad (4.8)$$

und

$$s_F = k_F \cdot F_{N_k(N_i-1)+N_k(N_j-1), N_k(N_i+N_j-1), 1-\alpha} \quad (4.9)$$

wobei $\chi_{N_k, 1-\alpha}$ das $1 - \alpha$ Quantil χ^2 -Verteilung mit N_k Freiheitsgraden und $F_{N_k(N_i-1)+N_k(N_j-1), N_k(N_i+N_j-1), 1-\alpha}$ das $1 - \alpha$ Quantil der F-Verteilung mit $N_k(N_i - 1) + N_k(N_j - 1)$ und $N_k(N_i + N_j - 1)$ Freiheitsgraden ist. Anstatt die schwer interpretierbaren Schwellwerte s_D und s_F direkt zu wählen, gibt der Nutzer die Parameter k_D , k_F und α vor, wobei α das Signifikanzniveau ist. k_D und k_F sind einfacher zu interpretieren, da sie an der Verteilung angepasst sind also auf statistische Größen zurückgeführt werden können.

Somit kann bei der Anwendung der Regeln, die in diesem Absatz beschrieben wurden, ein Satz von Tupeln von Regionen S_i und S_j erstellt werden. Diese werden sortiert nach D_{ij} . Das erste Element korrespondiert mit den zwei Regionen, die die höchste Ähnlichkeit der Grauwertverteilung besitzen und nicht durch eine signifikante Kante getrennt sind. Diese Regionen werden verschmolzen, inklusive der Grenzpixel, die sie bisher trennten. Der RAG wird anschließend aktualisiert. In diesem Zusammenhang müssen das Labelbild angepasst, die Attribute D_{ij} , F_{ij} und T_{ij} der verschmolzenen Region zu den benachbarten Regionen berechnet sowie die Kanten im RAG aktualisiert werden. Die Analyse und das Verschmelzen werden iterativ wiederholt, bis keine Regionen mehr verschmolzen werden können. Abbildung 12 zeigt das so entstandene Labelbild des Originalbildes auf der linken Seite der Abbildung.

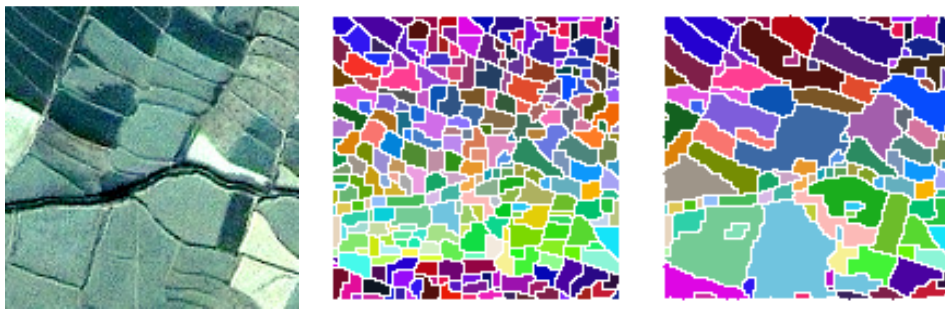


Abbildung 12: Links: Ackerlandobjekte aus Abbildung 11; Mitte: Bild nach der initialen Segmentierung; Rechts: erzeugtes Labelbild mit Hilfe des hier vorgestellten Segmentierungsverfahrens.

In Tabelle 3 werden die Parameter, die für den Segmentierungsprozess benötigt werden, zusammengefasst. Die Parameter der Segmentierung sind durch empirische Tests auf das verwendete Bildmaterial einzustellen, wobei ein einheitlicher Satz von Parametern für eine Szene (z.B. IKONOS-Szene) verwendet werden kann und eine weitere Adaptierung der Parameter für jedes GIS-Objekt nicht notwendig ist.

Größe	Beschreibung
σ	Parameter der Gaußfunktion, der den Grad der Glättung vor dem Anwenden des Wasserscheidenverfahrens bestimmt
α	Signifikanzniveau $\in [0,100\%]$;
s_T	Schwellwert für den maximalen Wert von T_{ij} mit $s_T \in [0,100\%]$; Stärke der Kante zwischen zwei Segmenten
k_D	Parameter für die Bestimmung des Schwellwert s_D nach Gleichung (4.8) für den maximalen Wert von D_{ij} aus Gleichung (4.3); Ähnlichkeit der Grauwertvektoren zweier Regionen
k_F	Parameter für die Bestimmung des Schwellwert s_F nach Gleichung (4.9) für den maximalen Wert von F_{ij} aus Gleichung (4.6)

Tabelle 3: Parameter für die Segmentierung.

4.4.3. Nachbearbeitung des Segmentierungsergebnisses an Hand von GIS-Spezifikationen

Wie in der Einleitung erwähnt, sind in vielen GIS kleine Teilflächen anderer Objektklassen innerhalb eines GIS-Objektes möglich. Die Existenz solcher kleinen Flächen einer anderen Objektklasse wird toleriert, solange diese kleiner als die im Objektartenkatalog vorgegebene Mindestkartierfläche sind. Diese Eigenschaft kann genutzt werden, um die Segmentierungsergebnisse nachzubearbeiten. Eine Region, die von nur einer weiteren Region umrandet wird und gleichzeitig die Mindestkartierfläche nicht überschreitet, wird mit dem umliegenden Segment verschmolzen, auch wenn die Bedingungen zum Verschmelzen dieser Regionen nicht erfüllt sind.

Ein Beispiel ist in Abbildung 13 gegeben. Das mittlere Bild, das das Segmentierungsergebnis eines aus zwei Schlägen bestehenden Ackerlandobjektes zeigt, weist eine Vielzahl kleiner Segmente auf (durch Kreise gekennzeichnet). Gründe hierfür sind u.a. Baumbepflanzungen am Rand des GIS-Objektes oder kleine Regionen, die sich auf Grund von Bodeneigenschaften signifikant von dem umliegenden Acker abheben und daher nicht verschmolzen wurden. Diese kleinen Segmente werden mit den umliegenden Segmenten verschmolzen, was zu einer Glättung des Segmentierungsergebnisses führt.

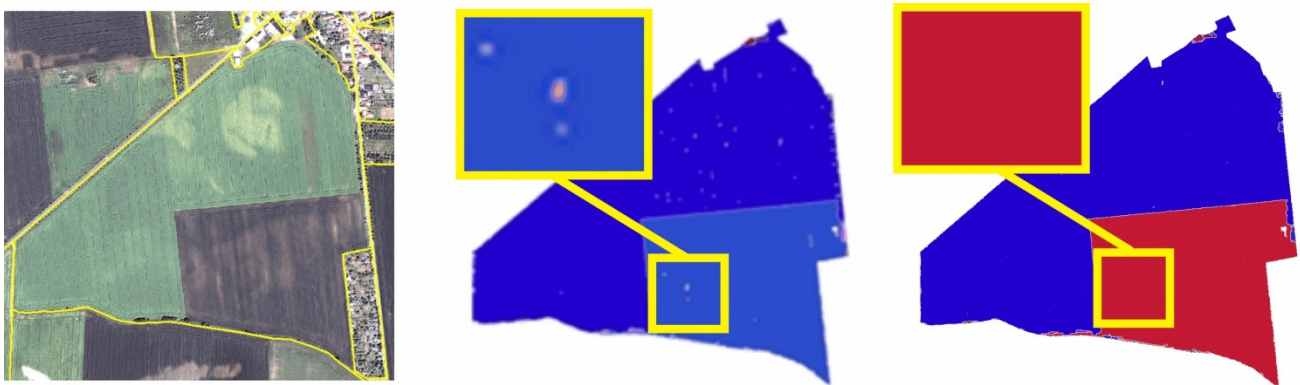


Abbildung 13: Verschmelzen kleiner Segmente, die nur ein Nachgarsegment besitzen (RGB- Bild eines GIS-Ackerlandobjektes mit zwei Ackerlandschlägen (links), Segmentierungsergebnis ohne Verschmelzung kleiner Objekte (Mitte) und Segmentierungsergebnis mit Verschmelzung kleiner Objekte (rechts)).

4.5. Klassifikation eines Segmentes

Nachdem die Segmente innerhalb der GIS-Objekte, die im ersten Analyseschritt als Acker-/Grünlandobjekte klassifiziert wurden, bestimmt sind, erfolgt nun die Klassifikation dieser Segmente in Acker- bzw. Grünland.

Die weiteren Analyseschritte werden auf den inneren Bereich der Segmente beschränkt. Dies ist für die Analyse notwendig, da es am Rand von Ackerland- und Grünlandflächen häufig zu Unregelmäßigkeiten kommen kann, die die Analyse der Objekte beeinflussen und das Ergebnis verfälschen können. Solche Unregelmäßigkeiten können Grünstreifen, Baumreihen, Teile einer Straße oder eines Weges, die in das Objekt ragen oder Spuren, die von wendenden landwirtschaftlichen Maschinen hinterlassen wurden, sein. Für die Beschränkung der Analyse auf das Innere wird auf jedes Segment ein morphologischer Filter (Erosion) angewendet, der die Seitenränder des Segmentes reduziert. Die Größe des verwendeten Filters hängt vom zu untersuchenden Gebiet und den dort vorkommenden landwirtschaftlichen Praktiken ab. Er kann einheitlich für eine zu untersuchende Szene genutzt werden.

Die Klassifikation eines Segmentes erfolgt anschließend in zwei Schritten. Zunächst werden für jedes Segment Merkmale bestimmt und diese in einem Merkmalsvektor gespeichert. An Hand dieses Merkmalsvektors wird anschließend das Segment mit Hilfe einer SVM klassifiziert.

4.5.1. Berechnung der Merkmale für die Klassifikation

Für die Klassifikation werden Merkmale aus den Merkmalsgruppen der spektralen, textuellen, strukturellen und geometrischen Merkmale benutzt. Da eine objektbasierte Klassifikation durchgeführt wird, werden alle Merkmale, bezogen auf ein gesamtes Ackerland-/Grünlandsegment ermittelt. Für Segmente, die eine vom Benutzer definierte Größe unterschreiten (z.B. die Mindestkartierfläche), werden keine Merkmale berechnet. Diese Segmente werden bei der anschließenden Bewertung einer Zurückweisungsklasse zugeordnet.

Die Erscheinungsform der Objektklasse Ackerland kann sehr stark variieren (z.B. Abbildung 16 und Abbildung 17). Sie hängt zum einen davon ab, ob der Acker überhaupt bewachsen ist und falls, mit welcher Fruchtart⁴. Ein bewachsenes Ackerlandobjekt ist ein Objekt, auf dem Vegetation vorhanden ist; ein unbewachsenes Ackerlandobjekt ist ein Objekt, auf dem keine Vegetation vorhanden ist. Der Unterschied zwischen nicht bewachsenem (z.B. Abbildung 16) und bewachsenem Ackerland (z.B. Abbildung 17) ist signifikant. Zur Erarbeitung der Strategie werden daher nicht die Objektklassen *Acker-* und *Grünland*, sondern die Objektklassen *Grünland*, *bewachsenes* und *nicht bewachsenes Ackerland* unterschieden. Die Unterscheidung von bewachsenen und unbewachsenen Ackerland erfolgt, um Variationen innerhalb der Objektklasse Ackerland zu verdeutlichen. Während der Klassifikation mittels SVM werden die Unterklassen nicht bewachsenes und bewachsenes Ackerland wieder zu einer Klasse zusammengefasst. Dies stellt für die Klassifikation mittels SVM kein Problem dar, weil die SVM in der Lage ist, Klassen zu trennen, auch wenn diese aus räumlich getrennten Clustern im Merkmalsraum bestehen. Wichtig ist, dass beim Training der SVM die unterschiedlichen Erscheinungsformen des Ackerlandes berücksichtigt werden und für jede Erscheinungsform genügend Trainingsgebiete zur Verfügung stehen.

Das Modell für die Unterscheidung von Ackerland (bewachsen und nicht bewachsen) und Grünland ist in Abbildung 14 in Form eines semantischen Netzes gegeben. Die Darstellung in einem semantischen Netz dient nur der Verdeutlichung dafür, dass eine Trennung von Acker- und Grünland mittels spektraler, textureller, struktureller sowie geometrischer Merkmale möglich ist. Die Klassifikation selbst wird aber nicht wissensbasiert, sondern wie oben erwähnt mit einer SVM durchgeführt. In Abbildung 14 ist zu sehen, dass die unterschiedlichen Klassen verschiedene spektrale, textuelle und strukturelle sowie Formmerkmale (geometrische Merkmale) besitzen. So kann nicht bewachsenes Ackerland von bewachsenem Ackerland auf Grund unterschiedlicher spektraler, textureller und struktureller Merkmale unterschieden werden. Bewachsenes Ackerland von Grünland kann auf Grund unterschiedlicher textureller, struktureller und Formmerkmale, und nicht bewachsenes Ackerland von Grünland auf Grund spektraler, struktureller und Formmerkmale unterschieden werden. In diesem Abschnitt sollen die verwendeten Merkmale konkretisiert werden, d.h. die spezifischen Merkmale, die in den Merkmalsvektor für die Klassifikation eingehen, werden vorgestellt und ihre Bestimmung detailliert beschrieben.

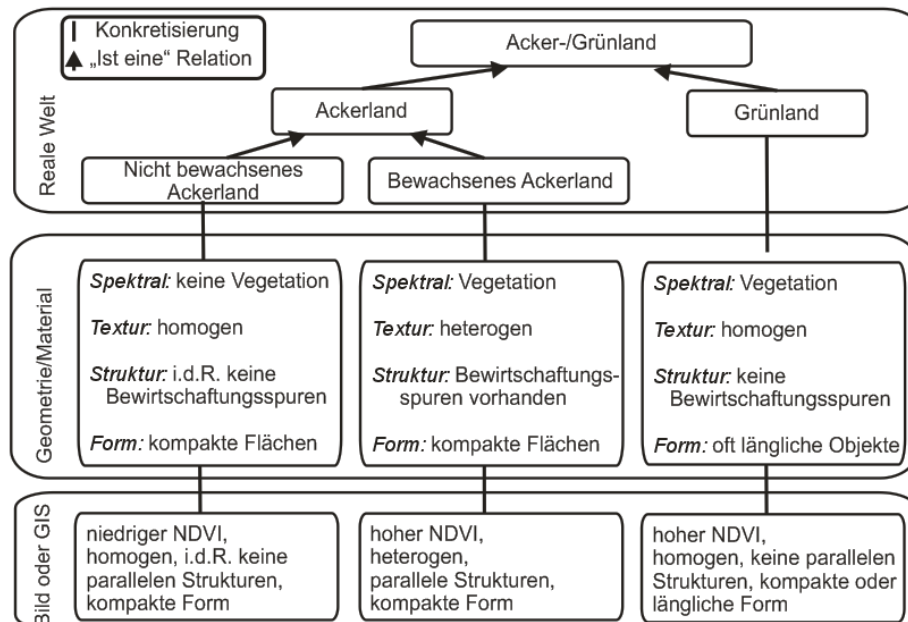


Abbildung 14: Modell zur Trennung von Acker- und Grünland mittels spektraler, textureller und struktureller Merkmale.

4.5.1.1. Spektrale Merkmale

Wie bereits im Kapitel 2 gezeigt, ist an Hand unterschiedlicher Reflexionseigenschaften die Trennung von landwirtschaftlichen Objektklassen, wie Acker- und Grünland, in Bildern möglich, was exemplarisch an Hand der Abbildung 15 bis Abbildung 17 zu sehen ist. Diese Abbildungen zeigen einen Ausschnitt eines

⁴ Eine Unterscheidung zwischen verschiedenen Fruchtarten ist nicht Ziel der Arbeit und wird nicht näher betrachtet.

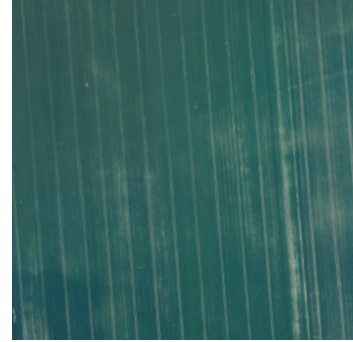
Grünlandobjektes sowie eines nicht bewachsenen und eines bewachsenen Ackerlandobjektes. Auch wenn die Vegetation auf dem Grünlandobjekt sehr spärlich ausfällt, ist diese klar zu erkennen, z.B. verglichen mit dem Bildausschnitt des nicht bewachsenen Ackerlandes in Abbildung 16. Abbildung 17 zeigt ein bewachsenes Ackerlandobjekt, das durch ein besonders saftiges Grün gekennzeichnet ist. Die Farbinformation der Bilder gibt also Hinweise darauf, ob Vegetation auf einem Objekt vorhanden ist oder nicht und ist somit ein Hinweis für die Objektklasse eines Segmentes. Jedoch kann mit Hilfe der spektralen Merkmale alleine nicht unterschieden werden, ob ein „saftiges“ Grünland- oder ein „saftiges“ bewachsenes Ackerlandsegment vorliegt. Dieses Problem kann zu bestimmten Zeitpunkten der Vegetationsperiode von Grün- und Ackerland vorliegen. Aus diesem Grund sind weitere Merkmale für die Klassifikation notwendig.



Abbildung 15: Grünland
(Luftbild 0,5m Bodenauflösung,
Mitte/Ende April).



**Abbildung 16: Nicht
bewachsenes Ackerland**
(Luftbild 0,5 m Bodenauflösung,
Mitte/Ende April).



**Abbildung 17: Bewachsenes
Ackerland (Luftbild 0,5 m
Bodenauflösung, Mitte/Ende
April).**

Die Farbinformationen aus den vorhandenen multispektralen Bildern werden direkt und indirekt zur Ableitung von Merkmalen für die Klassifikation verwendet. Bei einer direkten Ableitung werden Merkmale direkt aus den vom Sensor gespeicherten Daten, also auf die Werte der Farbbänder, erzeugt. Bei der indirekten Verwendung werden Merkmale aus Funktionen der Bänder berechnet. In dieser Arbeit wird, wenn möglich, der NDVI genutzt, der sich besonders gut für die Trennung von Flächen mit und ohne Vegetation eignet (Lillesand und Kiefer, 2000). Der NDVI berechnet sich aus den Kanälen *Nahes Infrarot* (NIR) und *Rot* wie folgt:

$$NDVI = \frac{NIR - Rot}{NIR + Rot} \quad (4.10)$$

Merkmale, die für die Klassifikation genutzt werden, sind somit:

- Mittelwerte (*avg*) pro Segment von jedem Band (und des NDVI)
- Standardabweichungen (*std*) pro Segment von jedem Band (und des NDVI)

Wenn nur ein RGB-Bild zur Verfügung steht und der NDVI nicht berechenbar ist (s. Gleichung (4.10)) ergibt sich daraus folgender Merkmalsvektor für die spektralen Merkmale:

$$\mathbf{x}_{spe_RGB} = (R_{avg}, G_{avg}, B_{avg}, R_{std}, G_{std}, B_{std})^T \quad (4.11)$$

Wenn ein Falschfarbenbild (IrRG) zur Verfügung steht, ergibt sich ein Merkmalsvektor für die spektralen Merkmale mit:

$$\mathbf{x}_{spe_IrRG} = (NIR_{avg}, R_{avg}, G_{avg}, NDVI_{avg}, NIR_{std}, R_{std}, G_{std}, NDVI_{std})^T \quad (4.12)$$

Und wenn ein multispektrales Bild mit den Farbkanälen *NIR*, *R*, *G* und *B* zur Verfügung steht, ergibt sich daraus folgender Merkmalsvektor für die spektralen Merkmale:

$$\mathbf{x}_{spe_IrRGB} = (NIR_{avg}, R_{avg}, G_{avg}, B_{avg}, NDVI_{avg}, NIR_{std}, R_{std}, G_{std}, B_{std}, NDVI_{std})^T \quad (4.13)$$

4.5.1.2. *Texturelle Merkmale*

Auch die Textur gibt Hinweise darauf, ob Grünland bzw. bewachsenes oder nicht bewachsenes Ackerland vorliegt. Während das Bild eines Grünlandsegmentes und eines nicht bewachsenen Ackerlandobjektes sehr homogen erscheinen (keine wiederkehrenden Muster innerhalb einer lokalen Nachbarschaft), erscheint das

Bild eines bewachsenen Ackerlandobjektes heterogener (wiederkehrende Muster), wie man exemplarisch Abbildung 15 bis Abbildung 17 entnehmen kann. Die Textur bietet daher die Möglichkeit, ein bewachsenes Ackerlandsegment (heterogen) von einem Grünland- oder unbewachsenen Ackerlandsegment (homogen) zu trennen.

Für die Analyse der textuellen Merkmale wird in dieser Arbeit die GLCM nach (Haralick et al., 1973) verwendet, die, wie im Kapitel 2 gezeigt, für die Klassifikation landwirtschaftlicher Flächen bereits erfolgreich eingesetzt wurde.

Die GLCM $P_{\Delta,\alpha}(g_1, g_2)$ beschreibt das gemeinsame Vorkommen von Grauwerten (g_1, g_2) zwischen Pixeln mit Abstand Δ und Winkel α und kann wie folgt ausgedrückt werden (Tönnies, 2005):

$$P_{\Delta,\alpha}(g_1, g_2) = \frac{1}{K} \sum_{\mathbf{x} \in R} \delta_D(g(\mathbf{x}) - g_1) \cdot \delta_D(g(\mathbf{x} + \mathbf{d}) - g_2) \quad (4.14)$$

mit

$$\mathbf{d} = \Delta \cdot (\sin \alpha \quad \cos \alpha)^T, \quad (4.15)$$

K als die Anzahl der Pixel eines Segmentes und δ_D als die Dirac-Funktion, wobei $\delta_D(0) = 1$ und $\delta_D(x \neq 0) = 0$ ist. In der Literatur (und auch in dieser Arbeit) erfolgt die Betrachtung über den Abstand von einem Pixel $\Delta = 1$ in vier unabhängige Richtungen der Achternachbarschaft, also über $\alpha = 0^\circ, 45^\circ, 90^\circ$ und 135° . Daraus ergeben sich vier GLCM: $P_{0^\circ}, P_{45^\circ}, P_{90^\circ}$ und P_{135° , die anschließend gemittelt werden ($P_{\Delta,\alpha} = (P_{0^\circ} + P_{45^\circ} + P_{90^\circ} + P_{135^\circ})/4$). Detaillierte Angaben zur Ableitung der GLCM sind (Haralick et al., 1973) und (Tönnies, 2005) zu entnehmen. Für die Erstellung der GLCM in dieser Arbeit wird als Eingangsbild, falls vorhanden, das Falschfarbenbild verwendet, alternativ wird das RGB-Bild verwendet. Das verwendete Eingabebild wird zunächst in ein Grauwertbild (*grau*) umgewandelt. Dabei wird folgende Formel verwendet (Russ, 1995):

$$\text{grau} = 0,299 \cdot \text{NIR} + 0,587 \cdot R + 0,114 \cdot G \quad (4.16)$$

bzw.

$$\text{grau} = 0,299 \cdot R + 0,587 \cdot G + 0,114 \cdot B \quad (4.17)$$

Aus der GLCM können dann Maße abgeleitet werden, die die Textur einer Region näher beschreiben. Die am häufigsten verwendeten Maße, die auch in dieser Arbeit verwendet werden, sind Energie, Kontrast, Homogenität und Korrelation. Die folgenden Erklärungen sind (Hall-Beyer, 2008) und die Formeln (Tönnies, 2005) entnommen.

Die Energie E (auch *Angular Second Moment* (ASM) genannt) ist ein Maß für die Gleichmäßigkeit der verschiedenen auftretenden Grauwertpaare. Die Energie ist hoch, wenn die Grauwertverteilung konstant oder periodisch ist. Sie berechnet sich aus der GLCM mit:

$$E = \sum_{g_1=0}^{K-1} \sum_{g_2=0}^{K-1} P_{\Delta,\alpha}^2(g_1, g_2) \quad (4.18)$$

Der Kontrast Kon beschreibt die mittlere Grauwertvariation und berechnet sich aus:

$$Kon = \sum_{g_1=0}^{K-1} \sum_{g_2=0}^{K-1} (g_1 - g_2)^2 \cdot P_{\Delta,\alpha}(g_1, g_2) \quad (4.19)$$

Die Homogenität H beschreibt die Ähnlichkeit von Nachbapixeln und berechnet sich aus:

$$H = \sum_{g_1=0}^{K-1} \sum_{g_2=0}^{K-1} \frac{1}{1+(g_1-g_2)^2} \cdot P_{\Delta,\alpha}(g_1, g_2) \quad (4.20)$$

Die Korrelation ϱ ist ein Maß für die lineare Abhängigkeit der Grauwerte mit benachbarten Pixeln. Sind regelmäßig wiederkehrende Muster (Grauwertpaare) vorhanden, ist die Korrelation hoch. Die Korrelation berechnet sich aus:

$$\varrho = \sum_{g_1=0}^{K-1} \sum_{g_2=0}^{K-1} \frac{P_{\Delta,\alpha}(g_1, g_2) - \mu_1 \mu_2}{\sigma_1^2 \sigma_2^2} \quad (4.21)$$

mit den Erwartungswerten μ_1 und μ_2 berechnet aus:

$$\mu_1 = \frac{1}{K} \sum_{g_1=0}^{K-1} g_1 \sum_{g_2=0}^{K-1} P_{\Delta,\alpha}(g_1, g_2) \quad (4.22)$$

$$\mu_2 = \frac{1}{K} \sum_{g_2=0}^{K-1} g_2 \sum_{g_1=0}^{K-1} P_{\Delta,\alpha}(g_1, g_2) \quad (4.23)$$

und dem Quadrat der Standardabweichungen σ_1 und σ_2 berechnet aus

$$\sigma_1^2 = \sum_{g_1=0}^{K-1} (g_1 - \mu_1)^2 \sum_{g_2=0}^{K-1} P_{\Delta,\alpha}(g_1, g_2) \quad (4.24)$$

$$\sigma_2^2 = \sum_{g_2=0}^{K-1} (g_2 - \mu_2)^2 \sum_{g_1=0}^{K-1} P_{\Delta,\alpha}(g_1, g_2) \quad (4.25)$$

Alle vier beschriebenen Haralick-Merkmale werden in einen Merkmalsvektor zusammengefasst:

$$\mathbf{x}_{\text{tex}} = (E, Kon, H, \varrho)^T, \quad (4.26)$$

4.5.1.3. Strukturelle Merkmale

Neben spektralen und textuellen Merkmalen wurden auch strukturelle Merkmale bereits für die Klassifikation landwirtschaftlicher Merkmale genutzt (s. Kapitel 2). Während die Textur das gemeinsame Vorkommen von Grauwerten beschreibt, beschreiben strukturelle Merkmale spezielle Strukturen in einem Segment, die mittels unterschiedlicher Techniken erkannt werden können und darauf basierend eine Klassifikation des Segmentes ermöglichen.

Ein Hauptunterscheidungsmerkmal im Erscheinungsbild von Acker- und Grünland liegt in dem Vorhandensein von Bearbeitungsspuren, die häufiger in der Klasse Ackerland auftreten. Diese Bearbeitungsspuren werden von landwirtschaftlichen Maschinen hinterlassen und verlaufen meist geradlinig parallel. Auch Feldfrüchte, wie z.B. Kartoffel- oder Erdbeerpflanzen, sind oft in geradlinig parallelen Reihen angepflanzt, die ebenfalls als geradlinige parallele Strukturen in Bildern erscheinen. Bearbeitungsspuren der landwirtschaftlichen Maschinen und Strukturen auf Grund der Feldfrucht werden in dieser Arbeit unter dem Begriff Bewirtschaftungsstrukturen zusammengefasst. Auf bewachsenem Ackerland können diese Bewirtschaftungsspuren i.d.R. beobachtet werden, wohingegen sie in unbewachsenem Ackerland weniger häufig auftreten. Auch dort sind Spuren nach der Saat vorhanden, die aber auf Grund des schlechten Kontrastes nicht immer im Bild sichtbar sind. Wenn diese geradlinig parallelen (Bewirtschaftungs-) Strukturen in einem Segment detektiert werden können, dann können diese Merkmale für eine Abgrenzung der Klasse bewachsenes Ackerland (geradlinig parallele Struktur vorhanden) gegenüber Grünland (keine geradlinig parallele Struktur vorhanden) genutzt werden. Eine Abgrenzung der Klasse unbewachsenes Ackerland gegenüber Grünland ist mittels spektraler Merkmale möglich, wie im Abschnitt 4.5.1.1 beschrieben. Strukturelle Merkmale wurden bisher vor allem für die Klassifikation von Weingärten und Plantagen genutzt. Wie in Kapitel 2 dargelegt, können diese Verfahren zur Extraktion struktureller Merkmale nicht ohne weiteres auf Ackerlandobjekte übertragen werden.

Der im folgenden Abschnitt vorgestellte Algorithmus zum Detektieren dieser geradlinig parallelen Spuren sowie das Ableiten von Merkmalen für die Klassifikation ist speziell auf die Problematik der Trennung von Acker- und Grünland abgestimmt. Da der Ansatz neuartig ist, soll dieser detailliert beschrieben werden. Der Algorithmus gliedert sich in vier Schritte: Vorverarbeitung, Erstellung eines Kantenbildes, Ableiten eines Richtungshistogramms aus dem Kantenbild und Ableitung von Merkmalen aus dem Richtungshistogramm für die Klassifikation. Alle Verarbeitungsschritte werden an Hand von vier Beispielobjekten mit Abbildungen und Zahlenwerten näher erläutert. Als Beispiele dienen neben einem Grünland- und einem nicht bewachsenen Ackerlandobjekt zwei bewachsene Ackerlandobjekte (Abbildung 18). Bei dem ersten bewachsenen Ackerlandobjekt sind die Bewirtschaftungsspuren gut zu erkennen; beim zweiten bewachsenen Ackerlandobjekt ist dies hingegen schwieriger. An Hand dieses Beispiels soll gezeigt werden, dass der Ansatz auch unter erschwerten Bedingungen in der Lage ist, die Klassen Grünland und bewachsenes Ackerland mit Hilfe struktureller Merkmale zu unterscheiden.

Vorverarbeitung

Zunächst sind mehrere Vorverarbeitungsschritte notwendig mit dem Ziel, die Strukturen von Interesse im Bild hervorzuheben. Dazu wird zunächst ein Intensitätsbild unter der Benutzung aller N_k Kanäle berechnet. Auf diesem Bild wird dann nach einer Histogrammlinearisation (auch Histogrammäqualisation oder Histogrammeinebnung genannt, engl. *Histogram Equalisation*) zur Verbesserung des Kontrastes eine anisotrope Filterung durchgeführt. Die anisotrope Filterung (Perona und Malik, 1990) betont besonders Strukturen, die im Bild vorhanden sind und ist daher für diese Anwendung besonders geeignet. Die Intensitätsbilder von Segmenten nach der Vorverarbeitung sind in Abbildung 18 (linke Spalte) zu sehen.

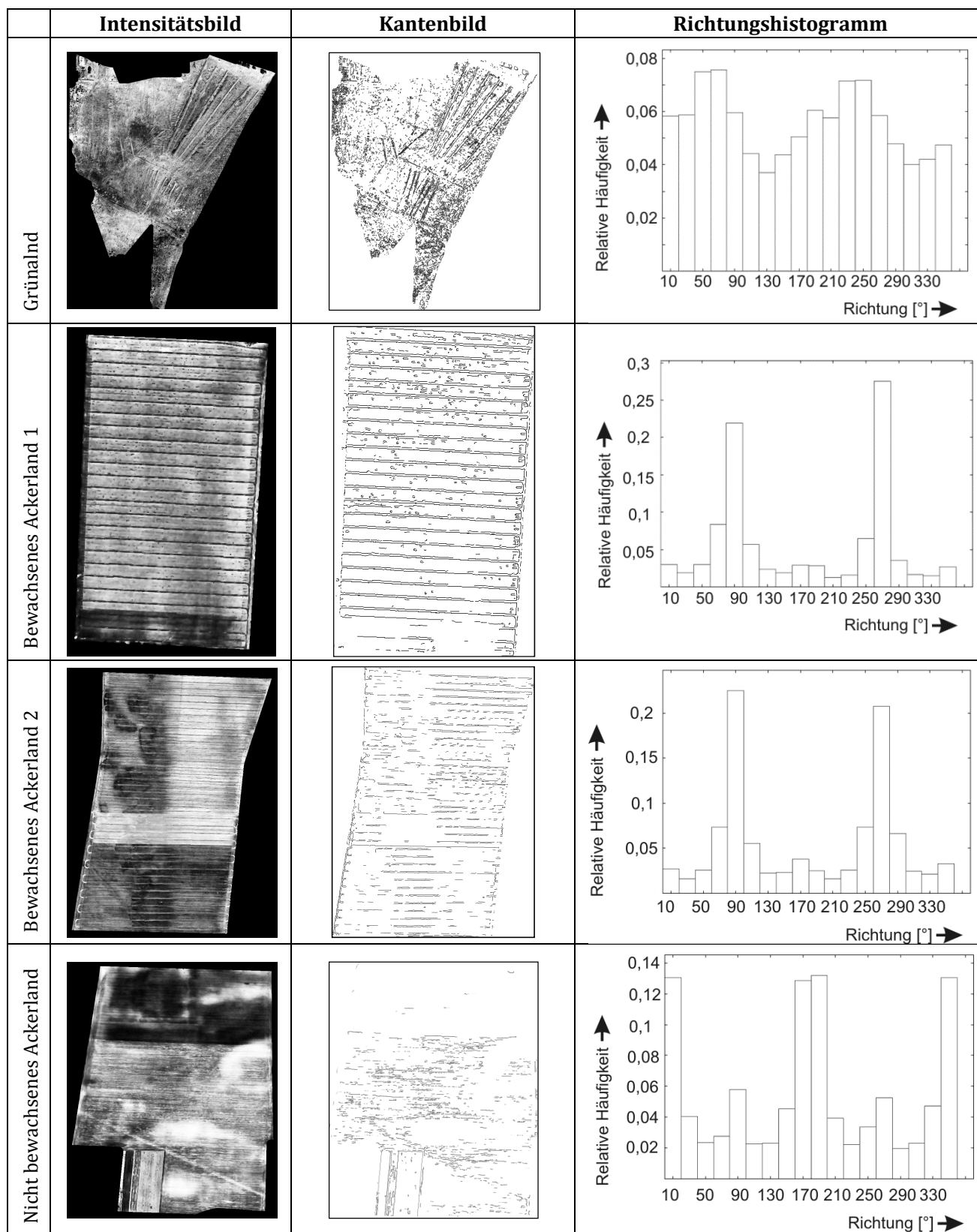


Abbildung 18: Beispiele für die Detektion und Merkmalsextraktion struktureller Merkmale an Hand von einem Grünlandobjekt (oben), zweier bewachsener (Mitte – die Intensitäts- und Kantenbilder beider Objekte sind um 90° gedreht) und einem nicht bewachsenen Ackerlandobjekt (unten) mit Intensitätsbild (links), Kantenbild (Mitte) und Richtungshistogramm (rechts). Die Histogramme besitzen unterschiedliche Skalen für die relative Häufigkeit!

Kantenbildextraktion

Die geradlinigen parallelen Bewirtschaftungsspuren, die es zu detektieren gilt, stellen sich in den Intensitätsbilder als Grauwertkanten dar. Daher wird nach der Vorverarbeitung ein Kantenfilter auf das Intensitätsbilder der Segmente angewendet. Untersuchungen haben gezeigt, dass sich für diese spezielle Anwendung der Canny-Operator (Canny, 1986) eignet. Das Ergebnis nach der Anwendung des Canny-Operators ist ein Kantenbild. Für jedes Kantenpixel stehen Amplitude und Richtung zur Verfügung. Durch empirische Tests wurden für den Canny-Operator der untere Hysteresisschwellwert auf 50 und der obere auf 70 festgelegt und diese Werte für alle Bilder gleichermaßen benutzt. Beispiele für ein Kantenbild sind in Abbildung 18 (mittlere Spalte) zu sehen. Die Richtungen können einen Wert von 0 bis 360° annehmen.

Histogrammerstellung

Mit Hilfe der Richtungs- und Amplitudenwerte der Kantenpixel soll nun geprüft werden, ob im Segment geradlinig parallele Strukturen vorhanden sind. Dazu wird ein Histogramm der Richtungen $h(r)$ aufgestellt, wobei die Klassenbreite des Histogramms 20° beträgt. Die Klassenbreite wurde auf 20° festgelegt, da in der Richtung der Bewirtschaftungsspuren leichte Variationen auftreten können, was unter anderem auf Unebenheiten im Boden, aber auch Ungenauigkeiten bei der Ausrichtung der Bewirtschaftung zurückgeführt werden kann. Diese Abweichungen sind i.d.R. nicht größer als 20°. Die Klassenbreiten werden nicht größer als 20° gewählt, da bei einer größeren Klassenbreite die Gefahr besteht, dass eine Differenzierung in Segmente mit oder ohne parallele Strukturen nicht mehr möglich ist. Bei der Erstellung wird für jedes Kantenpixel mit der Richtung r im Richtungshistogramm für diese Richtung eine Eintragung vorgenommen, wobei diese mit der Stärke der Kante (Amplitudenwert) gewichtet wird. Dadurch werden Kanten, die sich durch einen hohen Amplitudenwert auszeichnen, gegenüber schwächeren Kanten betont.

Nach der Erstellung des Histogramms sind weitere Schritte notwendig, bevor die Merkmale für die Klassifikation aus dem Histogramm abgeleitet werden können. So hat z.B. die Flächengröße des Segmentes Einfluss auf die Anzahl der gefundenen Kanten. Daher wird das Histogramm normalisiert, so dass das Histogramm die relative Häufigkeit für das Auftreten einer Richtung im Bild zeigt.

Ableiten der Merkmale

Beim Vergleichen der Richtungshistogramme in Abbildung 18 fällt zum einen auf, dass, während die Häufigkeiten der Richtungen im Grünland auf einem mehr oder weniger gleichen niedrigen Niveau liegen, dies bei den Ackerlandobjekten nicht der Fall ist. Da eine Kante mit der Richtung r_k dieselbe Ausrichtung wie die Kante r_{k+180° hat, besitzt Ackerland zwei signifikante Richtungen, die um 180° voneinander getrennt liegen. Diese Eigenschaft ist sogar bei dem nicht bewachsenen Ackerland zu erkennen. Dort sind schwache geradlinig parallele Linienstrukturen im Kantenbild sichtbar, die offensichtlich vom Pflügen herrühren.

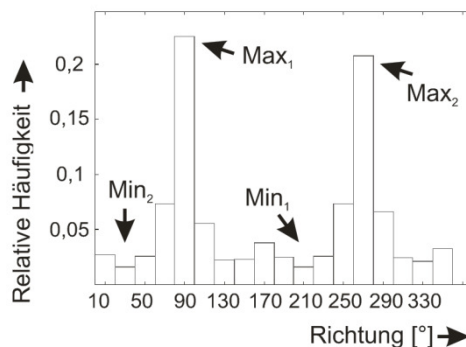


Abbildung 19: Einige Merkmale zur Klassifikation, abgeleitet aus dem Richtungshistogramm.

Die Unterschiede in den Histogrammen von Grünland gegenüber den Histogrammen von bewachsenem und nicht bewachsenem Ackerland lassen sich an Hand der folgenden Werte beschreiben. Sie sind, soweit möglich, in Abbildung 19 zur Erklärung dargestellt:

- kleinstes Minimum (Min_1),
- zweitkleinstes Minimum (Min_2),
- größtes Maximum (Max_1),
- zweitgrößtes Maximum (Max_2),
- Quotient q_1 aus Min_1 und Max_1 ($q_1 = Min_1/Max_1$),
- Quotient q_2 aus Min_1 und Max_2 ($q_2 = Min_1/Max_2$),

- Quotient q_3 aus Min_2 und Max_1 ($q_3 = Min_2/Max_1$),
- Quotient q_4 aus Min_2 und Max_2 ($q_4 = Min_2/Max_2$),
- Verhältnis der beiden größten Maximalwerte ($q_5 = 1 - Max_2/Max_1$),
- Standardabweichung der Werte des Histogramms σ_{rh} .

Einige dieser Merkmale sind korreliert. Auf eine Analyse zur Bestimmung der nicht korrelierten Merkmale wird in dieser Arbeit verzichtet. Dadurch, dass korrelierte Merkmale vorliegen, erhalten manche Merkmale ein stärkeres Gewicht. Die stärkere Gewichtung hat bei der Bestimmung, ob ein Segment mit parallelen Strukturen vorliegt, sogar einen positiven Effekt auf die Klassifikation, da diese Merkmale einen stärkeren Einfluss erhalten.

Damit ergibt sich ein Merkmalsvektor für die strukturellen Merkmale von:

$$\mathbf{x}_{\text{str}} = (Min_1, Min_2, Max_1, Max_2, q_1, q_2, q_3, q_4, q_5, \sigma_{rh})^T \quad (4.27)$$

Diese Größen sind für die Segmente aus Abbildung 18 in Tabelle 4 zusammengefasst. Tabelle 4 und den Histogrammen in Abbildung 18 ist zu entnehmen, dass die Minimalwerte für das Grünlandsegment immer größer und die Maximalwerte immer kleiner als die der Ackerlandsegmente sind. Aus dieser Beobachtung ergibt sich auch, dass die Quotienten der zuvor untersuchten Merkmale, sich für das Grünlandsegment deutlich gegenüber allen Ackerlandsegmenten unterscheiden. Ähnlich verhält es sich bei der Standardabweichung. Bei dem Maß $q_1 = 1 - Max_2/Max_1$ lässt sich hingegen an den gewählten Beispielen kein deutlicher Unterschied zwischen den einzelnen Segmenten erkennen.

Merkmals	Grünland	Bewachsenes Ackerland 1	Bewachsenes Ackerland 2	Unbewachsenes Ackerland
Min_1	0,04	0,01	0,03	0,02
Min_2	0,04	0,01	0,03	0,02
Max_1	0,08	0,28	0,12	0,13
Max_2	0,07	0,22	0,12	0,13
$q_1 = Min_1/Max_1$	0,49	0,05	0,25	0,15
$q_2 = Min_1/Max_2$	0,52	0,06	0,26	0,15
$q_3 = Min_2/Max_1$	0,53	0,05	0,28	0,17
$q_4 = Min_2/Max_2$	0,56	0,07	0,28	0,17
$q_5 = 1 - Max_2/Max_1$	0,05	0,20	0,08	0,01
σ_{rh}	0,01	0,07	0,06	0,04

Tabelle 4: Übersicht der strukturellen Merkmale der Objekte aus Abbildung 18.

4.5.1.4. Geometrische Merkmale

Es gibt eine Vielzahl von weiteren Merkmalen, die bisher in der Literatur als Zusatzinformation bei der Klassifikation von Acker- und Grünland genutzt wurden. Dazu gehören auch geometrische Merkmale (Formmerkmale). Grünlandobjekte treten häufig entlang von Straßen und Eisenbahnwegen auf und sind dadurch von einer schmalen und länglichen Form geprägt. Dies ist in Abbildung 20 zu sehen, die ATKIS GIS-Objekte der Klassen Ackerland (braun) und Grünland (grün) zeigt. In dieser Abbildung ist aber auch zu sehen, dass dies nicht immer der Fall sein muss (rechte untere und rechte obere Ecke). Informationen bezüglich der Form können daher, müssen aber nicht, eine Trennung von Acker- und Grünland ermöglichen.

Ein häufig verwendetes Formmerkmale ist die Kompaktheit k , die sich aus Umfang U und Fläche A eines Segmentes wie folgt berechnet:

$$k = \frac{A}{U}. \quad (4.28)$$

Die Kompaktheit bei schmalen und länglichen Segmenten, wie Grünlandsegmenten entlang von Straßen, Bächen und Schienenanlagen, ist geringer als bei quadratischen Segmenten, die häufiger der Klasse Ackerland angehören.

Weitere geometrische Formmerkmale können aus der Dimension einer begrenzende Hülle (*bounding box*) abgeleitet werden, wie sie in Abbildung 21b zu sehen ist (Tönnies, 2005). Ist ein Segment sehr schmal, dann ist l_2 klein, ist ein Objekt sehr lang, dann ist der Quotient q_6 mit

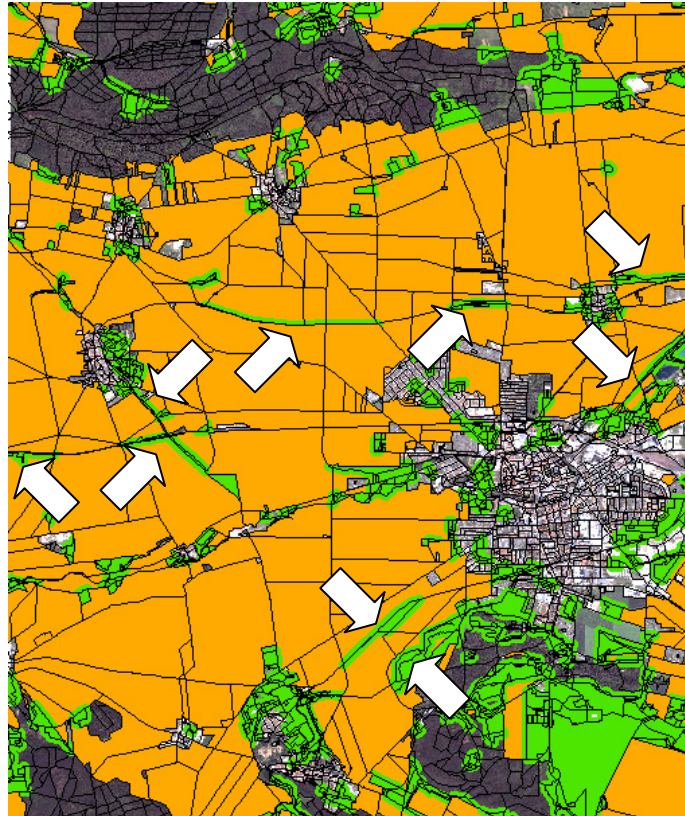


Abbildung 20: ATKIS GIS-Acker- (braun) und Grünlandobjekte (grün) für das Gebiet Halberstadt, Sachsen-Anhalt.

$$q_6 = l_2 / l_1 \quad (4.29)$$

klein. Beide Eigenschaften treffen oft, wie oben beschrieben, auf Grünlandsegmente zu.

Wenn das Segment jedoch sehr lang und schmal ist, aber viele Kurven aufweist, wie z.B. ein Grünlandsegment entlang eines Bachlaufes, dann kann dies nicht aus l_2 und q_6 abgeleitet werden. Ein Maß, das das Verhältnis der Fläche der bounding box A_{bb} mit der Fläche des Segments A vergleicht, ist in diesem Fall besser geeignet. Dieses Maß wird Konvexität kon genannt und berechnet sich aus (Burger und Burge, 2006):

$$kon = A / A_{bb}. \quad (4.30)$$

Ist die Konvexität klein, gilt oben beschriebener Fall, ist die Konvexität groß, dann liegt ein Segment wie in Abbildung 21 vor.

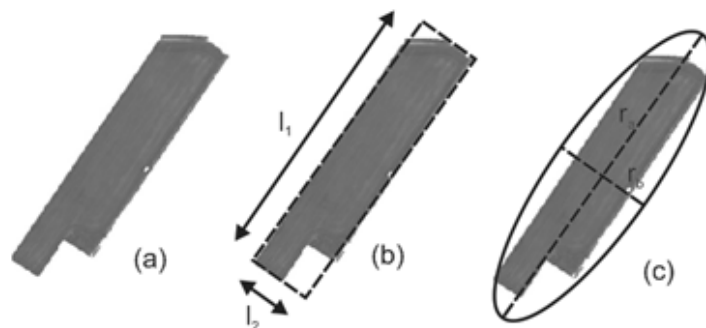


Abbildung 21: Bounding Box (b) und umschließende Ellipse (c) eines Segmentes (a).

Ein Merkmal, das ebenfalls wie q_6 in (4.29) die längliche Ausdehnung eines Segmentes prüft, ist die Exzentrizität e , die sich aus der Größe der Halbachsen r_a und r_b einer Ellipse, die ein Segment umschließt (Abbildung 21c), berechnet. Die Formel dafür lautet (Burger und Burge, 2006):

$$e = \frac{(r_a - r_b)}{r_a}. \quad (4.31)$$

Handelt es sich um ein rundes Segment, dann beträgt die Exzentrizität Null, ein sehr schmales und langgezogenes Segment hat einen Exzentrizitätswert nahe eins.

Weitere Formmerkmale, die sich aus r_a und r_b einer umschließenden Ellipse ableiten lassen, sind u.a. die Anisometrie (*Anisometry*) a , die Sperrigkeit (*bulkiness*) b und der Strukturfaktor sF , die wie folgt berechnet werden (Burger und Burge, 2006):

$$a = \frac{r_a}{r_b}, \quad (4.32)$$

$$b = \pi \frac{r_a r_b}{A} = \frac{A_{ell}}{A}, \quad (4.33)$$

$$sF = a \cdot b - 1, \quad (4.34)$$

wobei A_{ell} der Flächeninhalt der umschließenden Ellipse ist.

Zusammenfassend kann gesagt werden: Die Formmerkmale, die in dieser Arbeit verwendet werden, sind:

- Kompaktheit k eines Segmentes (4.28),
- Wert der kleineren Seite der bounding box l_2 ,
- Quotient aus den Seitenlängen der bounding box q_6 (4.29),
- Konvexität kon eines Segmentes (4.30),
- Exzentrizität e eines Segmentes (4.31),
- Anisometrie a eines Segmentes (4.32),
- Sperrigkeit b eines Segmentes (4.33),
- Strukturfaktor sF eines Segmentes (4.34).

Auch hier liegen ähnlich wie bei den strukturellen Merkmalen korrelierte Merkmale vor. Wie im Abschnitt 4.5.1.3 bereits diskutiert, wird auch bei den geometrischen Merkmalen auf eine Bestimmung der nicht korrelierten Merkmale verzichtet.

Damit ergibt sich ein Merkmalsvektor für die geometrischen Merkmale von:

$$\mathbf{x}_{geom} = (k, l_2, q_6, kon, e, a, b, sF)^T \quad (4.35)$$

4.5.1.5. Zusammenfassung

Für die Klassifikation eines Segmentes wird ein Merkmalsvektor $\mathbf{x}$ gebildet, in der Form:

$$\mathbf{x} = (\mathbf{x}_{spe}^T, \mathbf{x}_{tex}^T, \mathbf{x}_{stru}^T, \mathbf{x}_{form}^T)^T \quad (4.36)$$

wobei je nach zur Verfügung stehenden Bilddaten $\mathbf{x}_{spe} = \mathbf{x}_{spe_RGB}$ bzw. $\mathbf{x}_{spe} = \mathbf{x}_{spe_IrRG}$ oder $\mathbf{x}_{spe} = \mathbf{x}_{spe_IrRGB}$ ist. Die Dimension des Merkmalsvektors $\mathbf{x}$ schwankt daher zwischen 28 und 32 Merkmalen (s. Tabelle 5).

Der so berechnete Merkmalsvektor $\mathbf{x}$ eines Segmentes dient als Eingabe in die Klassifikation mittels SVM.

Bilddaten	Dimension $\mathbf{x}$
RGB	28
IrRG	30
RGB und NIR	32

Tabelle 5: Dimension des Merkmalsvektors in Abhängigkeit von den zur Verfügung stehenden Bilddaten.

4.5.2. Klassifikation anhand der Merkmale

Für die Klassifikation mittels SVM werden die in Abschnitt 4.5.1 eingeführten Unterklassen (nicht bewachsenes Ackerland und bewachsenes Ackerland) wieder zu einer Klasse Ackerland zusammengefasst. Die Klassen, die getrennt werden sollen, sind also *Acker-* und *Grünland*.

Zunächst erfolgt das Training. Beim Training werden entweder vom menschlichen Bearbeiter repräsentative GIS-Objekte oder zufällig gewählte GIS-Objekte des zu verifizierenden GIS ausgewählt. Wichtig ist, wie bereits erwähnt, dass beim Training der SVM die unterschiedlichen Erscheinungsformen des Ackerlandes berücksichtigt werden und für jede Erscheinungsform genügend Trainingsgebiete zur Verfügung stehen. Die für das Training gewählten GIS-Objekte sollten nur einen Schlag enthalten und besitzen eine korrekte Klassenzugehörigkeit im GIS. Für alle Trainings-GIS-Objekte werden die Merkmale berechnet und so der jeweilige Merkmalsvektor $\mathbf{x}_i$ bestimmt (s. Abschnitt 4.5.1). Bevor diese für das Anlernen der SVM benutzt werden, werden alle Merkmale linear in den Bereich von -1 und +1 skaliert (s. Abschnitt 3.1.3). Die Skalierungswerte werden für die spätere Klassifikation von Segmenten, deren Klassenzugehörigkeiten unbekannt sind, gespeichert.

Beim Training der SVM werden neben der Lage der Hyperebene im Merkmalsraum auch die Parameter C und γ mittels Gittersuche unter Verwendung der Kreuzvalidierung automatisch trainiert (s. Abschnitt 3.1.4). Es wurde die RBF-Kernfunktion verwendet, da sie sich in der Literaturstudie in Abschnitt 3.1.1 als gut geeignet herausgestellt hat.

Bei der Klassifikation eines Segmentes wird ebenfalls zunächst der Merkmalsvektor berechnet (s. Abschnitt 4.5.1) und dieser anschließend mit den gespeicherten Skalierungswerten aus dem Training skaliert. Der skalierte Merkmalsvektor wird dann der SVM übergeben, die mittels Lage gegenüber der Hyperebene das Segment der Klasse Ackerland oder Grünland zuweist.

4.6. Klassifikation und Verifikation des GIS Objektes

Abschließend muss die Klassifikation und die Verifikation eines GIS-Objektes vorgenommen werden. Bei der Klassifikation eines GIS-Objektes werden zuerst alle Segmente des GIS-Objektes einzeln klassifiziert. Unterschreitet ein Segment die vom Benutzer vorgegebene Mindestfläche, wird dieses Segment einer *Zurückweisungsklasse* zugeschlagen. Ist die Flächensumme der *Zurückweisungsklasse* größer als die Flächensumme der Segmente, die klassifiziert werden konnten, gilt das Objekt als nicht klassifizierbar. Die klassifizierten Segmente sind entweder der Klasse *Ackerland* oder *Grünland* mittels Klassifikation zugewiesen wurden. Das GIS-Objekt wird der Klasse zugeordnet, deren Flächensumme die Mindestkartierfläche überschreitet. Überschreiten die Flächensummen beider Klassen die Mindestkartierfläche bzw. wird von keiner Klasse die Mindestkartierfläche erreicht, wird das GIS-Objekt der *Zurückweisungsklasse* zugeordnet.

Im Verifikationsschritt erfolgt der Vergleich des Klassifikationsergebnisses des GIS-Objektes mit der Information in der GIS-Datenbank. Wurde eine andere Klasse (oder die *Zurückweisungsklasse*), als in der GIS-Datenbank gespeichert, bestimmt, wird dieses GIS-Objekt vom System als falsch zurückgewiesen. Ergebnis des Verifikationsprozesses ist die Aussage für jedes GIS-Objekt, ob dieses vom System als falsch bewertet oder angenommen wurde. Alle als falsch bewerteten Objekte werden dem menschlichen Bearbeiter zur visuellen Kontrolle vorgelegt.

4.7. Nachbearbeitung

Die Nachbearbeitung besteht aus der Verknüpfung der Verifikationsergebnisse mit dem GIS-Datensatz und der anschließenden Visualisierung. Ein Beispiel für eine Visualisierung eines verifizierten GIS ist in Abbildung 22 gegeben. Die vom System akzeptierten Objekte, die in der Abbildung grün gekennzeichnet sind, brauchen vom Menschen nicht mehr näher betrachtet zu werden. Dadurch entsteht für den Bearbeiter eine Zeitersparnis. Nur die vom System als falsch bewerteten Objekte, die in der Abbildung rot gekennzeichnet sind, müssen vom Bearbeiter kontrolliert werden.

4.8. Grenzen des Verfahrens

Bisher wurden keine Angaben über Grenzen des Algorithmus gemacht. Das soll in diesem Abschnitt erfolgen. Eine wesentliche Voraussetzung für alle Verarbeitungsschritte ist natürlich, dass die Objekte im Bild nicht verdeckt sind (z.B. durch Schnee oder Wolken).

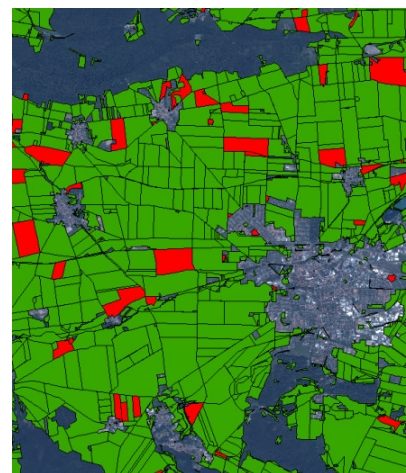


Abbildung 22: Beispiel eines Verifikationsergebnisses von GIS-Objekten der Klasse Acker- und Grünland (grün – vom System akzeptierte Objekte; rot – vom System als falsch bewertete Objekte).

4.8.1. Pixelbasierte Klassifikation zur Identifizierung von Ackerland- / Grünlandobjekten

Erfahrungen zeigen, dass nicht bewachsene kleine Ackerlandobjekte auf Grund ähnlicher spektraler (niedriger NDVI) und textueller Eigenschaften (homogene kleine kompakte Gebiete) mit *Industrie* verwechselt werden können. Zwei solcher Beispiele sind in Abbildung 23 dargestellt. Diese Charakteristik führt zur Ablehnung des GIS-Ackerlandobjektes. Der menschliche Bearbeiter muss sich dieses GIS-Objekt, zusätzlich zu allen anderen abgelehnten Objekten, anschauen. Fehler verbleiben damit nicht im GIS, nur der manuelle Aufwand der Nachbearbeitung steigt.



Abbildung 23: In Cyan hervorgehoben sind kleine nicht bewachsene GIS-Ackerlandobjekte, die als Industrie (gelb) falsch klassifiziert wurden. Datengrundlage ist ein multispektrales IKONOS-Bild und ATKIS. Klassifikationsergebnisse: Siedlung: rot, Wald: grün, Ackerland/Grünland: braun, Industrie: gelb).

4.8.2. Segmentierung von Schlägen

Ziel der Segmentierung ist das Finden homogener Flächen. Wenn die Fläche eines Schlages innerhalb eines GIS-Objektes nicht homogen ist, wie dies bei unterschiedlicher Bodenfeuchtigkeit oder dem Auftreten eines anderen Bodentyps entstehen kann, dann kann es zu einer Übersegmentierung kommen. Ein Beispiel hierfür ist in Abbildung 24 gegeben, dort liegt in dem rechten unteren Bereich des GIS-Objektes im rechten Schlag eine andere Bodenbeschaffenheit vor, obwohl es mit der gleichen Feldfrucht bewachsen ist.

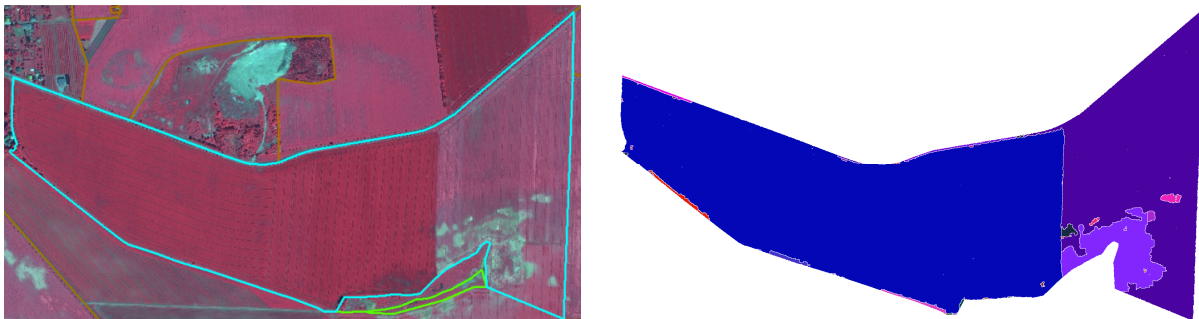


Abbildung 24: GIS-Ackerlandobjekt mit unterschiedlichen Bodenbedingungen (rechts unten) und 2 Schlägen mit unterschiedlicher Bewirtschaftungsrichtung, aber gleichem homogenen Erscheinungsbild (links).

Ein anderes Beispiel für eine Übersegmentierung ist gegeben, wenn nicht homogene Flächen wie Bäume oder Plantagen in einem GIS-Objekt vorkommen. Im GIS-Objekt in Abbildung 25 sind Bäume an der Objektgrenze zu erkennen.



Abbildung 25: GIS-Grünlandobjekt mit Bäumen am Objekttrand (Dimension des Ausschnittes ca. 300 x 300 m).

Die vorgestellten Fälle, die alle eine Übersegmentierung des Bildes zur Folge haben, sind für den Algorithmus nicht kritisch, da z.B. die Summe der Flächen der Bäume in Abbildung 25 nicht den Schwellwert für die Zulässigkeit dieser Objektart in einem GIS-Ackerlandobjekt überschreitet. Dieser Schwellwert ist Bestandteil des Algorithmus und im jeweiligen GIS-Objektartenkatalog definiert. Sollte die Baumreihe jedoch den Schwellwert überschreiten und daher einen kompakten Fehler darstellen, so sollte dieser Objektteil bereits vorher vom Klassifikationsalgorithmus zur Abgrenzung von Ackerland- und Grünlandobjekten zu anderen Objektklassen detektiert werden.

Ein Beispiel für eine Untersegmentierung ist in der Abbildung 24 zu sehen, bei dem zwei Schläge mit sehr ähnlichen spektralen Eigenschaften, aber einer unterschiedlichen Hauptbewirtschaftungsrichtung im linken Teil des GIS-Objektes verschmolzen wurden. Eine Untersegmentierung kann dazu führen, dass das Segment, das zwei Schläge enthält, einer falschen Klasse zugeordnet wird bzw. Schläge, die nicht in einer bestimmten Objektklasse zugelassen sind und die Mindestkartierfläche überschreiten, nicht detektiert werden können. Dies kann dazu führen, dass entweder dem menschlichen Bearbeiter ein Mehraufwand entsteht, da dieses GIS-Objekt vom System abgelehnt wird, oder dass ein Fehler im GIS verbleibt, der vom System nicht gefunden werden kann.

4.8.3. Klassifikation eines Segmentes

Auf die Nachteile und Grenzen des benutzten Klassifikationsalgorithmus, der SVM, wurde bereits im Kapitel 3 eingegangen. Der Fokus der Diskussion der Grenzen liegt daher auf den für die Klassifikation verwendeten Merkmalen.

4.8.3.1. Spektrale Merkmale

Häufig wird in der Diskussion über die Verwendung von spektralen Eigenschaften bei einer objektbasierten Klassifikation der Einfluss von Mischpixeln als Nachteil genannt. Bei den Bilddaten, für die der in dieser Arbeit vorgestellte Algorithmus ausgelegt ist (geometrische Auflösung von 0,5 – 1 m), ist der Einfluss von Mischpixeln zu vernachlässigen, da Segmente, die so klein sind, dass eventuell ein Mischpixeleinfluss vorliegt, auf Grund ihrer Größe als nicht klassifizierbar eingestuft und somit der Zurückweisungsklasse zugeschrieben werden.

Bei einer geometrischen Auflösung von 0,5 – 1 m werden oft Schattenflächen als ein Problem angeführt. Diese Flächen sind bei den Klassen Ackerland und Grünland vernachlässigbar klein (z.B. bei Baumreihen am Feldrand) bzw. treten erst gar nicht auf.

4.8.3.2. Texturelle Merkmale

Die Berechnung der GLCM und damit der texturellen Merkmale hängt von zwei Parametern ab, nämlich dem Abstand Δ und dem Winkel α (s. Gleichung (4.15)). Die berechneten Texturmerkmale sind also nur repräsentativ für diese gewählten Parameter. Da nicht nur ein Winkel α , sondern vier Winkel genutzt und die Ergebnisse aus allen Richtungen gemittelt werden, ist die GLCM rotationsinvariant, die Abhängigkeit des Texturmaßes von α und der Richtungsabhängigkeit der Textur konnte so weitgehend eliminiert werden. Die Abhängigkeit gegenüber dem Abstand Δ kann nicht eliminiert werden, so dass nur Textureigenschaften, die innerhalb der Reichweite von Δ liegen, bewertet werden können. Texturen, die weiträumiger sind und nicht innerhalb von Δ liegen werden nicht berücksichtigt.

4.8.3.3. Strukturelle Merkmale

Der Denkansatz besteht in der Annahme, dass ein Grünlandsegment keine geradlinigen parallelen Bewirtschaftungsspuren enthält. Dies ist bis auf wenige Ausnahmen im Jahr auch gegeben. Eine Ausnahme ist, dass ein Grünlandsegment z.B. gemäht wird und dann ebenfalls geradlinige parallele Strukturen enthalten kann. In Mitteleuropa erfolgt das Abmähen der Wiesen 2- bis 3-mal im Jahr (Juni/Juli, August und September), dabei verbleibt das abgemähte Gras ca. eine Woche auf den Wiesen zum Trocknen. Findet eine Bilderfassung innerhalb dieser Woche statt, können die Grünlandsegmente, auf die der beschriebene Sachverhalt zutrifft, auf Grund der strukturellen Merkmale falsch klassifiziert werden.

Umgekehrt kann von der Annahme ausgegangen werden, dass in bewachsenem Ackerland keine Bearbeitungsspuren zu sehen sind. Falls keine Bearbeitungsspuren im Bild zu erkennen sind, u.a. weil sie tatsächlich nicht vorhanden sind oder die Auflösung des Bildes zu klein ist, um diese erkennen zu können, wird eine Trennung zwischen Acker- und Grünland ebenfalls schwer möglich sein.

Des Weiteren ist die Differenzierung (Acker-/Grünland) auf die Geradlinigkeit und Parallelität der Bewirtschaftungsspuren angewiesen, weil diese als Entscheidungsmerkmale zur Trennung von Acker- und Grünland dient. Nicht gerade oder nicht parallele Bewirtschaftungsspuren können aus unterschiedlichen Gründen entstehen. Dies kann der Fall sein, wenn zum einen die Wirtschaftlichkeit keine andere Form zulässt (Hanglage oder besonders kleine Schläge; z.B. Abbildung 26) oder wenn Kreisbewässerungen vorliegen (z.B. Abbildung 27). Der erste Fall tritt zumindest in Nord- und Mitteldeutschland eher selten auf, da der Anbau auf diesen Flächen unwirtschaftlich ist. Diese Flächen sind oft Grünlandflächen und enthalten dann selten ausgeprägte Strukturen (s. oben). Auch Kreisbewässerungen sind zumindest in Deutschland unüblich; sie sind eher in trockenen Regionen zu finden. Wenn der Algorithmus jedoch in Regionen mit Kreisbewässerung eingesetzt werden soll, so ist durch eine kleine Modifikation die Erkennung konzentrischer Kreise möglich. Dies erfolgt durch Verwendung eines geeigneten Akkumulationsraumes. Nach der Extraktion eines Kantenbildes wird dieses in einen geeigneten Akkumulationsraum überführt. Im Folgenden wird der Hough-Raum als Beispiel für einen Akkumulationsraum verwendet. Die Maxima im Hough-Raum repräsentieren die Kreismittelpunkte. Ist ein Ackerlandobjekt, wie in Abbildung 27 vorhanden, bilden sich signifikante Maxima heraus. Aus dem Hough-Raum lässt sich dann ein Histogramm mit der Häufigkeit der auftretenden Mittelpunkte ableiten, mit dem dann analog wie mit dem Richtungshistogramm in der vorliegenden Arbeit verfahren werden kann.

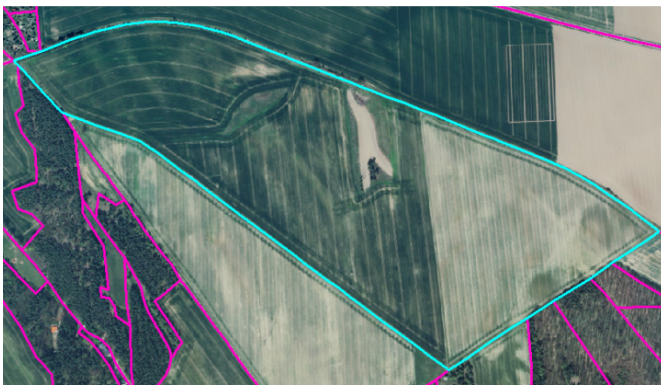


Abbildung 26: Nicht geradlinige Bewirtschaftung eines Ackerlandschlages (links oben).

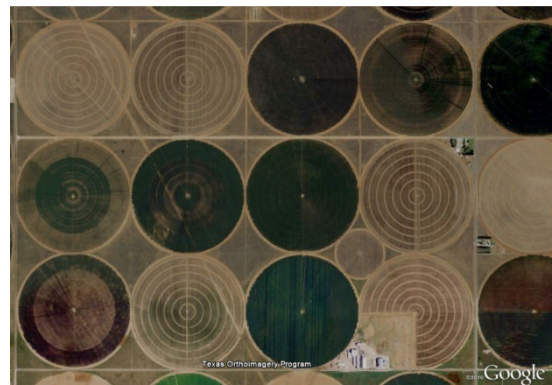


Abbildung 27: Kreisförmige Bewirtschaftung in Texas, USA (Quelle: Google).

Eine andere Einschränkung der strukturellen Merkmale ergibt sich über die Begrenzung, dass eine Mindestgröße für die Ackerlandfläche zum Detektieren der Spuren vorhanden sein muss. Das heißt, wenn Segmente zu klein (speziell zu schmal) werden, ist keine zuverlässige Analyse der Strukturen möglich. Dies ist besonders häufig bei Grünlandobjekten der Fall, da diese entlang von Gewässern, Straßen und Eisenbahnlinien auftreten, wie es in Abbildung 28 zu sehen ist. Es müssen daher für eine zuverlässige spektrale, texturelle und strukturelle Analyse Schwellwerte eingeführt werden, die die Form des Segmentes näher beschreiben. In diesem Fall ist die Breite des Segmentes (oder des GIS-Objektes) ebenfalls ein Kriterium. Ein im GIS geführtes Ackerlandobjekt, welches so schmal ist, dass kein Pflug mehr eingesetzt werden kann, ist wahrscheinlich ein Grünlandobjekt. Die Breite wird durch die geometrischen Merkmale berücksichtigt.

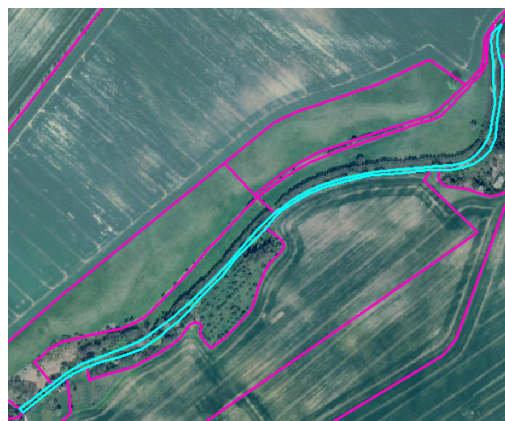


Abbildung 28: Beispiel für ein langes und schmales GIS-Grünlandobjekt in Cyan entlang eines Bachlaufes.

4.8.3.4. Geometrische Merkmale

Geometrische Merkmale bringen nur dann eine Zusatzinformation bei der Klassifikation, wenn deren Form repräsentativ ist. Kommt es z.B., wie in Abbildung 24 exemplarisch gezeigt, zu Fehlsegmentierungen, können geometrische Merkmale im Gegensatz zu spektralen, texturellen und strukturellen Merkmalen keine nutzbaren Information für die Klassifikation liefern.

4.9. Zusammenfassung

In diesem Kapitel wurde eine Methode vorgestellt, mit der erstmals die Verifikation von GIS-Objekten der Objektarten Acker- und Grünland möglich ist. Das beschriebene Verfahren ist in zwei Schritte aufgeteilt. Zunächst wird überprüft, ob ein GIS-Objekt der Klassen Acker- oder Grünland tatsächlich Acker- oder Grünland enthält. Dies geschieht mit Hilfe eines pixelbasierten Klassifikationsverfahrens (Gimel'farb, 1996) mit anschließender Übertragung des Ergebnisses auf das GIS-Objekt mit dem in der Praxis bereits erfolgreich getesteten Ansatzes von Busch et al. (2005). Anschließend werden die GIS-Objekte, die tatsächlich mit Acker- oder Grünland bedeckt sind, mittels eines weiteren Analyseschrittes untersucht, der dann das GIS-Objekt entweder der Klasse *Ackerland*, *Grünland* oder einer *Zurückweisungsklasse* zuordnet. Bei diesem Klassifikationsschritt werden Bedingungen, die durch den GIS-Objektartenkatalog gegeben sind (speziell für GIS, die den Inhalt und Detailierungsgrad einer topographischen Karte mittleren Maßstabs entsprechen) beachtet. Dazu zählt u.a. das Vorhandensein unterschiedlicher Schläge innerhalb eines GIS-Acker- oder Grünlandobjektes. Nachdem die Klassifikation des GIS-Objektes abgeschlossen ist, erfolgt die Verifikation durch Vergleich des Klassifikationsergebnisses mit der im GIS gegebenen Objektart. Das Ergebnis der Verifikation, die Entscheidung über Annahme oder Ablehnung eines GIS-Objektes, wird visualisiert. Die endgültige Kontrolle aller vom System abgelehnten GIS-Objekte erfolgt durch einen menschlichen Bearbeiter.

5. Evaluation

Ziel der Evaluation ist die Überprüfung der Verwendbarkeit der in Kapitel 4 vorgestellten zweistufigen Methode zur Verifikation von GIS-Ackerland- und Grünlandobjekten. Zunächst wird auf die Bewertung von und die Kriterien für die Verifikations- und Klassifikationsergebnisse eingegangen. Nachdem die verwendeten Bild- und GIS-Daten erläutert wurden, werden die benötigten Parameter sowie deren Einstellung detailliert beschrieben. Anschließend werden erst beide Analyseschritte getrennt und abschließend das gesamte Verfahren auf mehreren Testszenen evaluiert.

5.1. Bewertung der Kriterien für die Klassifikations- und Verifikationsergebnisse

5.1.1. Bewertung von Klassifikations- und Verifikationsergebnissen

Sowohl bei der Bewertung der Klassifikation als auch der Verifikation kommen Entscheidungsmatrizen (auch Konfusionsmatrizen nach dem englischen Wort *confusion matrices* genannt) zum Einsatz (Lillesand und Kiefer, 2000), (Congalton und Green, 2009). Ein Beispiel für eine Konfusionsmatrix zur Bewertung eines Klassifikationsergebnisses von 2 Klassen (Objektart 1 und Objektart 2) ist in Abbildung 29 zu sehen. Pro Klasse lassen sich die Genauigkeit aus Nutzersicht (*user's accuracy*) und die Herstellergenauigkeit (*producer's accuracy*) berechnen. Der Nutzer ist an der *Korrektheit* eines Klassifikationsergebnisses interessiert, also wie viele Pixel/Objekte einer Klasse im Klassifikationsergebnis tatsächlich dieser Klasse angehören. Die *user's accuracy* berechnet sich aus dem Quotienten der Anzahl der korrekt erkannten Pixel/Objekte einer Klasse mit der im Klassifikationsergebnis angegebenen Anzahl dieser Klasse:

$$user's\ accuracy_i = \frac{e_{ii}}{\sum_j e_{ij}} \cdot 100\%, \quad (5.1)$$

wobei e die Einträge in der Konfusionsmatrix sind und der erste Index die Objektart des Klassifikationsergebnisses und der zweite die Referenz wiedergibt (Abbildung 29). Der Hersteller hingegen ist daran interessiert, eine gegebene Klasse möglichst gut zu erkennen, also wie viele Pixel/Objekte der Referenz richtig erkannt werden (*Vollständigkeit*). Die *producer's accuracy* berechnet sich als Quotient aus der Anzahl richtig erkannten Pixel/Objekte einer Klasse und der Gesamtanzahl der Pixel/Objekte dieser Klasse, die in der Referenz enthalten sind:

$$producer's\ accuracy_i = \frac{e_{ii}}{\sum_i e_{ij}} \cdot 100\%, \quad (5.2)$$

wobei die Indices analog der *user's accuracy* vergeben sind. Während die *user's accuracy* von Fehlern 1. Art (*error of commission*) beeinflusst sind, ist die *producer's accuracy* von Fehlern 2. Art (*error of omission*) beeinflusst.

Klassifikationsergebnis \ Referenz	Objektart 1	Objektart 2
	Objektart 1	Objektart 2
Objektart 1	e_{11}	e_{12}
Objektart 2	e_{21}	e_{22}

Abbildung 29: Konfusionsmatrix zur Bewertung eines Klassifikationsergebnisses.

Weitere Parameter für die Bewertung eines Klassifikationsergebnisses sind die Gesamtgenauigkeit (*overall accuracy*) und der Cohen's Kappa-Index, der auch Kappa-Koeffizient genannt wird. Die *overall accuracy* ist der Quotient aus der Anzahl aller korrekten Zuordnungen und der Anzahl aller Pixel/Objekte, also

$$overall\ accuracy = \frac{\sum_i e_{ii}}{\sum_{ij} e_{ij}} \cdot 100\%. \quad (5.3)$$

Der Vorteil der *overall accuracy* ist, dass nur ein Maß für die Bewertung der Klassifikation angegeben wird, während die *user's accuracy* und *producer's accuracy* jeweils pro Klasse angegeben werden. Ein Nachteil besteht darin, dass keine Informationen enthalten sind, ob die Fehler gleich über die Klassen verteilt sind oder manche Klassen gut und andere schlecht erkannt werden. Dafür ist der Kappa-Koeffizient besser geeignet. Der Kappa-Koeffizient ist ein Maß für die Übereinstimmung von Klassifikationsergebnis und Referenz, das aus der Konfusionsmatrix bestimmt und wie folgt berechnet wird (Greve und Wentura, 1997):

$$\kappa = \frac{p_0 - p_c}{1 - p_c} \cdot 100\% \quad (5.4)$$

mit p_0 als tatsächliche Übereinstimmung berechnet aus den Diagonalelementen der Konfusionsmatrix mit

$$p_0 = \frac{\sum_i e_{ii}}{N}, \quad (5.5)$$

wobei N die Gesamtanzahl der Elemente der Konfusionsmatrix sowie p_c als erwartete Übereinstimmung, berechnet aus den Zeilen und Spaltensummen der Konfusionsmatrix mit

$$p_c = \frac{\sum_i (\sum_j e_{ji} \cdot \sum_j e_{ij})}{N^2}. \quad (5.6)$$

Ein negativer Kappa-Koeffizient weist darauf hin, dass die tatsächliche Übereinstimmung geringer als die Zufallserwartung ist. Nach Fleiss (1983) ist ein Kappa-Koeffizient zwischen 40% und 60% annehmbar, zwischen 60% und 75% als gut und darüber hinaus als ausgezeichnet anzusehen.

Bei der Bewertung von Verifikationsergebnissen kommen ebenfalls Konfusionsmatrizen zum Einsatz, wobei die vom Algorithmus verworfenen/akzeptierten Objekte/Pixel mit dem in der Referenz als korrekt/falsch gekennzeichneten Objekten/Pixel verglichen werden. Wenn die Referenz mit den zu verifizierenden Daten (z.B. GIS-Daten) übereinstimmt, so wird dieses Objekt/Pixel in der Referenz als korrekt geführt; gibt es Abweichungen, wird dieses Objekt/Pixel in der Referenz als falsch gekennzeichnet. Bei der Konfusionsmatrix für die Verifikationsergebnisse handelt sich also um eine spezielle Form bzw. Interpretation der zuvor vorgestellten Konfusionsmatrix zur Bewertung der Klassifikationsergebnisse. Bei der Konfusionsmatrix für die Verifikation unterscheidet man zwischen *True Positive (TP)*, *False Negative (FN)*, *False Positive (FP)* und *True Negative (TN)*, die sowohl absolut als auch in [%] angegeben werden können und in Abbildung 30 dargestellt sind. Die Summe aus *TP*, *FN*, *FP* und *TN* ergibt dabei 100%. Die Summe aus *FP* und *TN* entspricht dem im Ausgangsdatensatz enthaltenen Fehlern, die Summe aus *TP* und *FN* ist der Anteil der korrekten Objekte/Pixel.

Referenz \ BA-Operator	akzeptiert	abgelehnt
	akzeptiert	abgelehnt
korrekt	True Positive (<i>TP</i>)	False Negative (<i>FN</i>)
falsch	False Positive (<i>FP</i>)	True Negative (<i>TN</i>)

Abbildung 30: Konfusionsmatrix zur Bewertung eines Verifikationsergebnisses.

Der Erfolg eines Verifikationsergebnisses kann mittels aus der Konfusionsmatrix abgeleiteten Maße beurteilt werden. Ein Maß, das den Anteil der korrekten Objekte in der GIS-Datenbank wiedergibt, ist die *thematische Genauigkeit (TG)*. Diese wird unterschieden in die *TG*, die vor dem Verifikationsprozess (*TG a priori*) und nach dem Verifikationsprozess (*TG a posteriori*) vorlag. *TG a priori* und *TG a posteriori* lassen sich wie folgt berechnen:

$$TG \text{ a priori} = 100\% - (FP[\%] + TN[\%]) = TP[\%] + FN[\%], \quad (5.7)$$

$$TG \text{ a posteriori} = 100\% - FP[\%]. \quad (5.8)$$

Ist die *TG* hoch, ist der Anteil der im GIS enthaltenen Fehler gering. Die *TG* verschiedener Objektarten innerhalb eines GIS unterscheidet sich, da jede Objektart eine unterschiedliche Anzahl von Fehlern enthält. Ein weiteres Maß ist die Effizienz, die die Summe der Objekte/Pixel in [%] ist, die vom Verifikationsalgorithmus akzeptiert werden:

$$Effizienz = TP[\%] + FP[\%] \quad (5.9)$$

und damit beschreibt, wie viele Objekte bei einer manuellen Nachkontrolle nicht mehr betrachtet werden müssen (wirtschaftliche Effizienz). Die Berechnung der *overall accuracy* nach (5.3) für die Verifikation ergibt sich aus

$$overall \text{ accuracy} = TP[\%] + TN[\%]. \quad (5.10)$$

Die *overall accuracy* beschreibt also die Effizienz des Systems in Hinblick auf die korrekt bewerteten GIS-Objekte. Ein weiteres Kriterium ist die Fehlerdetektionsrate (*error detection rate*). Sie beschreibt den Anteil der gefundenen Fehler im Verhältnis der Fehler im GIS vor dem Verifikationsprozess und berechnet sich aus:

$$\text{error detection rate} = \frac{TN}{TN+FP} \cdot 100\%. \quad (5.11)$$

5.1.2. Kriterien für die Verifikationsergebnisse

Für die Verifikation müssen die Klassifikationsergebnisse so gut sein, dass eine zuverlässige Verifikation möglich ist. Wie bereits im ersten Kapitel beschrieben, ist die Überprüfung der *TG* Schwerpunkt dieser Arbeit. In der Ausschreibung zur Aktualisierung des Digitalen Landschaftsmodells Deutschland (DLM-DE) durch das Bundesamt für Kartographie und Geodäsie (BKG) ist eine *TG* von 95%, bezogen auf den Flächenanteil und die Anzahl der korrekten Objekte, gefordert (BKG, 2009). Das DLM-DE basiert auf dem ATKIS Basis-DLM. Das ATKIS Basis-DLM und somit auch das DLM-DE sind Vertreter eines GIS, die dem Inhalt und Detaillierungsgrad einer topographischen Karte mittleren Maßstabs entspricht, die auch im Fokus dieser Arbeit stehen. Das vom BKG genannte Ziel einer *TG* von 95%, bezogen auf den Flächenanteil und die Anzahl der Objekte, soll daher auch als Kriterium der Bewertung des in Kapitel 4 vorgestellten Verfahrens dienen. Nach dem Verifikationsprozess soll also die *TG a posteriori* der Klassen Acker- und Grünland mindestens die geforderten 95% erreichen. Außerdem soll einem menschlichen Bearbeiter bei der Anwendung des in dieser Arbeit vorgestellten semi-automatischen Systems eine Zeitersparnis bei der Qualitätskontrolle eines GIS gegenüber der rein manuellen Kontrolle entstehen. Daher soll die *Effizienz* so hoch wie möglich sein (Busch et al., 2005). Wenn sehr viele Fehler in einem GIS enthalten sind, ist die *Effizienz* nicht alleine geeignet, um eine Aussage über die Güte des Verifikationsprozesses zu treffen, da durch einen sehr hohen Anteil von *TN* die *Effizienz* gering ist. In diesen Fällen wird zusätzlich sowohl die *overall accuracy* als auch die *error detection rate* betrachtet, um eine Aussage über die Güte des Verifikationsergebnisses treffen zu können.

5.2. Datengrundlage

Zunächst werden die für die Tests zur Verfügung stehenden GIS und Bilddaten allgemein beschrieben, bevor die einzelnen Testgebiete im Detail vorgestellt werden.

5.2.1. GIS-Datensätze

Für die Evaluierung standen zwei verschiedene GIS zur Verfügung. Zum einen handelt es sich um ATKIS (speziell das ATKIS Basis-DLM), zum anderen um selbst erzeugte Schlagkataster. Für alle GIS-Datensätze liegen Referenzdaten vor, die manuell von einem versierten Bildauswerter auf der Grundlage der zur Verfügung stehenden Bilddaten erfasst wurden.

5.2.1.1. ATKIS

ATKIS ist ein bundesweites Projekt der Arbeitsgemeinschaft der Vermessungsverwaltungen der Länder der Bundesrepublik Deutschland (AdV). Mit ATKIS wird die Topographie der Bundesrepublik Deutschland in einer geotopographischen Datenbasis beschrieben und in Form digitaler Erdoberflächenmodelle bereitgestellt (AdV, 2011). Im Rahmen von ATKIS gibt es verschiedene digitale Erdoberflächenmodelle. Das Produkt mit der höchsten Auflösung, das Grundlage für alle weiteren digitalen Landschaftsmodelle ist, ist das Digitale Basis-Landschaftsmodell (Basis-DLM). Das Basis-DLM liegt flächendeckend für Deutschland vor und wird u.a. auf der Grundlage von Luftbildern mit einer Auflösung von 20 bis 40 cm, unterstützt von Vor-Ort-Erfassungen, aktualisiert. Damit wird eine Lagegenauigkeit von Straßen, Gewässern und Schienenwegen von ± 3 m gewährleistet. Der Zyklus der Aktualisierung von ATKIS liegt bei 5 Jahren, wobei Objekte von besonderer Bedeutung der Spitzenaktualität unterliegen und innerhalb von 3 (z.B. Straßen), 6 (z.B. Flughafen) und 12 (z.B. Schienenbahn) Monaten aktualisiert werden.

Beschreibungen, wie Objektartnummer der Klasse (OBJART) sowie deren Definition, der Objekttyp und die Mindestkartierfläche sind dem ATKIS Objektartenkatalog des Basis-DLM (AdV, 1997) zu entnehmen und für die Klassen Acker- und Grünland in Tabelle 6 zusammengefasst. Eine wesentliche Eigenschaft der ATKIS-GIS-Objekte der Klassen Acker- und Grünland ist, dass diese mehr als eine Bewirtschaftungseinheit (Schlag) beinhalten können und dies in der Regel auch tun. Für Grünland können zusätzlich Werte für das Attribut Funktion (FKT) angegeben werden (Landwirtschaftsfläche (2730) und Verkehrsbegleitgrün (2740)). Diese Werte sind im Normalfall jedoch nicht vergeben. Daher soll der in Kapitel 4 vorgestellte Algorithmus, der eigentlich für Grünlandobjekte der Unterklasse FKT 2730 entwickelt worden ist, auf alle im GIS als Grünland

geführten Objekte angewendet werden. Auf Grund von Generalisierung sind innerhalb eines GIS-Ackerland- oder Grünlandobjektes auch Baumreihen als Schlagbegrenzung oder versiegelte Flächen bei straßenbegleitendem Grünland möglich.

Objektart	OBJART	Definition	Objekttyp	Mindestkartierfläche
Ackerland	4101	Fläche für den Anbau von Feldfrüchten (z.B. Getreide, Hülsenfrüchte, Hackfrüchte) und Beerenfrüchten (z.B. Erdbeeren ⁵).	Flächenförmig	Fläche ≥ 1 ha
Grünland	4102	Gras- und Rasenflächen, die gemäht oder beweidet werden.	Flächenförmig	Fläche ≥ 1 ha

Tabelle 6: Definition von Acker- und Grünland nach ATKIS Objektartenkatalog (AdV, 1997).

Weiterhin ist anzumerken, dass, obwohl die Mindestkartierfläche sowohl für Ackerland als auch für Grünland 1 ha beträgt, diese Mindestkartierfläche in realen Datensätzen häufig unterschritten wird. Die Anzahl dieser GIS-Objekte ist exemplarisch für drei GIS-Datensätze in Tabelle 7 zusammengefasst.

Testszene	Objektart	Anzahl der Objekte kleiner als 1 ha	Prozent der Objekte kleiner als 1 ha
Halberstadt	Ackerland	24	7,32%
	Grünland	127	28,29%
Hildesheim	Ackerland	68	14,38%
	Grünland	49	21,21%
Weiterstadt	Ackerland	68	9,91%
	Grünland	49	24,62%

Tabelle 7: Anzahl der GIS-Objekte unter der Mindestkartierfläche.

Neben den Objektarten Acker- und Grünland sind in den ATKIS-Datensätzen vor allem auch die Objektarten Siedlung (2111 bzw. 2113), Industrie (2112 bzw. 2114) und Wald (4107) vorhanden, deren Spezifikationen in Tabelle 8 zusammengefasst sind. Auch hier kann es zu Generalisierungen kommen, so dass in einem Siedlungsobjekt Bäume und kleine Grünflächen auftreten können. Die Mindestkartierfläche von Wald im ATKIS beträgt 0,1 ha, Siedlung und Industrie sind immer zu erfassen, wobei aus Erfahrung gesagt werden kann, dass diese Flächen eine Mindestgröße von 0,1 ha besitzen.

Objektart	OBJART	Objekttyp	Mindestkartierfläche
Wohnbaufläche	2111	Flächenförmig	vollzählig
Fläche gemischter Nutzung	2113		
Industrie- und Gewerbefläche	2112	Flächenförmig	vollzählig
Fläche besonderer funktionaler Prägung	2114		
Wald, Forst	4107	Flächenförmig	0,1 ha

Tabelle 8: Definitionen anderer flächenmäßig häufiger auftretender GIS-Objektarten nach ATKIS Objektartenkatalog (AdV, 1997).

5.2.1.2. Schlagkataster

Neben ATKIS-Datensätzen standen für die Evaluierung auch Schlagkataster zur Verfügung. Die Schlagkataster wurden manuell erstellt und entsprechen den Vorgaben des Objektartenkataloges des ATKIS Basis-DLM mit zwei Ausnahmen. Zum einen darf ein GIS-Objekt der Klasse Ackerland oder Grünland eines Schlagkatasters im Unterschied zum ATKIS nur einen Schlag enthalten (daher auch der Name „Schlagkataster“). Zum anderen liegen die Mindestkartierflächen bei 0 ha. Es werden also alle Schläge einzeln erfasst.

5.2.2. Bilddaten

Für die Evaluierung liegen sowohl Satelliten- als auch Luftbilder vor. Alle verwendeten Bilddaten sind orthorektifiziert und multispektral mit den Kanälen *NIR*, *R*, *G* und *B*. Die geometrischen Auflösungen der Luft- und Satellitendaten unterscheiden sich und liegen bei 0,4 m bzw. 1 m. Bei den Satellitenbilddaten mit 1 m geometrischer Auflösung handelt es sich um pansharp IKONOS-Bilder. Die Luftbilder wurden aus einem kleinen Flugzeug heraus mit dem Niedrigpreissystem PFIFF (Grenzdörffer, 2005) aufgenommen. Der

⁵ Anmerkung der Autorin: Die Erdbeere zählt aus botanischer Sicht nicht zu den Beeren, sondern zu der Gruppe der Sammelnussfrüchte. Die Definitionen von Acker- und Grünland in Tabelle 6 wurden (AdV, 1997) entnommen.

Vorteil dieses Systems ist, dass es relativ unabhängig vom Wetter verwendet werden kann, da es auch bei Bewölkung einsetzbar ist. Auf Grund von Kalibrierungsproblemen sind die Luftbildinformationen im *NIR*-Kanal nur bedingt nutzbar. Ein erster Überblick über die Bilddaten ist in Tabelle 9 gegeben.

Bilder	Bänder	Geometrische Auflösung
IKONOS	<i>NIR, R, G, B</i>	1 m
PFIFF	<i>NIR, R, G, B</i>	0,4 m

Tabelle 9: Übersicht über die Bilddaten für die Evaluierung.

5.2.3. Testszenen

Für die Evaluierung standen die in Tabelle 10 zusammengefassten Testszenen zur Verfügung. Bei den Testgebieten handelt es sich um typische Szenen für Nord- und Mitteldeutschland. Zur Testszene Rukieten ist anzumerken, dass auch wenn ein multitemporaler Bilddatensatz zur Verfügung stand, die Bilddaten für jeden Zeitpunkt separat ausgewertet werden. Im Folgenden sollen die verwendeten Daten näher beschrieben werden.

Ort	GIS			Bilddaten	
	Schlagkataster	Ausgedünntes Schlagkataster	ATKIS	Sensor	Zeitpunkt der Aufnahme
Halberstadt	X	X	X	IKONOS	18. Juni 2005
Hildesheim			X	IKONOS	02. April 2005
Weiterstadt			X	IKONOS	24. Juni 2003
Rukieten	X			PFIFF	Verschieden (s. Tabelle 20)

Tabelle 10: Testszenen für die Evaluierung.

5.2.3.1. Halberstadt, Sachsen-Anhalt

Die Szene Halberstadt (kurz: HBS) ist ein repräsentatives Beispiel einer Region in Deutschland mit intensiver landwirtschaftlicher Nutzung (ca. 72% der Szene ist mit Acker- oder Grünland bedeckt), wobei auch große Wald- und Siedungsflächen sowie einige Industrieflächen in dieser Szene zu finden sind. An diesem Datensatz ist anzumerken, dass sich die GIS-Ackerlandobjekte durch besonders große Schläge auszeichnen. Für dieses Testgebiet stehen insgesamt drei GIS zur Verfügung – ein Schlagkataster, ATKIS und ein ausgedünntes Schlagkataster.

Schlagkataster

Im Schlagkataster der Testszene Halberstadt liegen neben Ackerland- und Grünlandobjekten auch Objekte der Klassen Siedlung (2111, 2113), Industrie (2112, 2114) und Wald (4107) vor. Statistiken zum Schlagkataster bezogen auf Ackerland- und Grünlandobjekte sind in Tabelle 11 und Tabelle 12 zusammengefasst. Während sich die Angaben in Tabelle 11 auf die Anzahl der GIS-Objekte beziehen, sind in Tabelle 12 die Flächengrößen angegeben. Die Daten in Tabelle 11 und Tabelle 12 beziehen sich nur auf Acker- und Grünland bzw. der gemeinsamen Klasse beider Objektarten, da nur diese mit Hilfe des in Kapitel 4 vorgestellten Verfahrens verifiziert werden sollen. In den Tabellen werden jeweils zunächst die Anzahl der Objekte/deren Fläche in der zweiten und anschließend die enthaltenen Fehler in der dritten Spalte aufgeführt. In der vierten Spalte wird angegeben, bei wie vielen von den in der dritten Spalte genannten Fehlern es sich um Verwechslungen zwischen den Klassen Acker- und Grünland oder umgekehrt handelt, die gleichzeitig größer als 0,5 ha sind (Mindestgröße eines Segmentes zur Durchführung des zweiten Analyseschrittes – s. Abschnitt 5.3). Bei dieser Art von Fehlern handelt es sich also genau um die Fehler, die mit dem zweiten Analyseschritt entdeckt werden sollen. Alle anderen Fehler werden in dieser Spalte nicht berücksichtigt. Da die Anzahl der Fehler somit geringer ist, unterscheiden sich die *TG a priori* je Objektklasse der Spalten drei und vier; die Werte in Spalte 4 fallen immer höher aus. In Klammern sind die jeweiligen *TG a priori* angegeben. Eine *TG a priori* von 100% bedeutet, dass keine Fehler in der GIS-Datenbank enthalten sind. Wie man sehen kann, liegen die *TG a priori* bezogen auf die Anzahl der GIS-Objekte immer unter den geforderten 95%, während zumindest bei Ackerland, bezogen auf die Fläche, die *TG a priori* sogar bei über 95% liegt. Ist die *TG* bezogen auf die Fläche höher als die *TG* bezogen auf die Anzahl der GIS-Objekte, spricht dies dafür, dass die fehlerhaften GIS-Objekte eher klein sind. In der letzten Spalte ist angegeben, wie viele der GIS-Objekte als Trainingsgebiete für die objektbasierte Klassifikation zufällig ausgewählt wurden (graphisch dargestellt in Abbildung 32).

Objektklasse	Anzahl d. Objekte	enthaltene Fehler	Verwechslung zwischen Acker- u. Grünland	Trainingsgebiete für die objektbasierte Klass.
Ackerland	493	69 (86,0%)	43 (88,4%)	150
Grünland	454	174 (61,7%)	43 (84,5%)	83
Acker-/Grünland	947	243 (74,3%)	86 (88,4%)	233
Andere Klassen	1699	2	-	
Gesamt	2646	245	86	

Tabelle 11: Statistiken über GIS-Ackerland- und Grünlandobjekte des Schlagkatasters der Testszene HBS, bezogen auf die Anzahl der Objekte. In Klammern ist die jeweilige *TG a priori* angegeben.

Objektklasse	Fläche d. Objekte	enthaltene Fehler	Verwechslung zwischen Acker- u. Grünland	Trainingsgebiete für die objektbasierte Klass.
Ackerland	85 km ²	3,10 km ² (96,4%)	1,80 km ² (97,9%)	28,75 km ²
Grünland	15 km ²	6,61 km ² (56,9%)	5,29 km ² (62,2%)	3,06 km ²
Acker-/Grünland	100 km ²	9,71 km ² (90,3%)	7,09 km ² (92,8%)	31,81 km ²
Andere Klassen	39 km ²	0,02 km ²	-	
Gesamt	139 km ²	9,73 km ²	7,09 km ²	

Tabelle 12: Statistiken über GIS-Ackerland- und Grünlandobjekte des Schlagkatasters der Testszene HBS, bezogen auf die Flächengröße. In Klammern ist die jeweilige *TG a priori* angegeben.

Für das Training des pixelbasierten Klassifikationsverfahrens des ersten Analyseschrittes zur Trennung der gemeinsamen Klasse Acker-/Grünland gegenüber anderen Klassen wie Siedlung, Industrie und Wald werden die in Abbildung 31 verwendeten Trainingsregionen genutzt. Während die Trainingsgebiete für die objektbasierte Klassifikationsverfahren zufällig ausgewählt wurden, werden die Trainingsgebiete der pixelbasierten Klassifikation vom Benutzer ausgewählt.

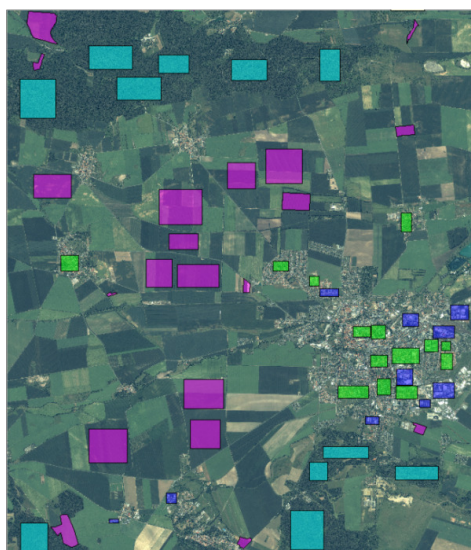


Abbildung 31: Trainingsgebiete pixelbasierte Klassifikation HBS (grün: Siedlung, blau: Industrie, magenta: Wald, violett: Acker-/Grünland).



Abbildung 32: Trainingsgebiete objektbasierte Klassifikation HBS (blau: Ackerland, magenta: Grünland).

ATKIS

Dieselben Statistiken wie für das Schlagkataster werden auch für ATKIS angegeben und sind in Tabelle 13 und Tabelle 14 zusammengefasst. Die Bedeutung der Spalten in diesen Tabellen ist identisch mit der in Tabelle 11 und Tabelle 12. Auch hier lässt sich beobachten, dass die *TG a priori*, bezogen auf die Anzahl der GIS-Objekte, immer unter 95% liegt und außerdem sehr viel geringer als die *TG a priori*, bezogen auf die Flächengröße, ist. Die *TG a priori* bezogen auf die Flächengröße liegt immer bei über 95%, Ausnahme ist die Klasse Grünland. Dies ist wiederum ein Hinweis darauf, dass die fehlerhaften GIS-Objekte eher kleine Objekte sind. Für die pixelbasierte und objektbasierte Klassifikation wurden dieselben Trainingsgebiete wie für das Schlagkataster der Testszene Halberstadt verwendet, die in Abbildung 31 und Abbildung 32 dargestellt sind. Für die objektbasierte Klassifikation ist dies notwendig, da der Algorithmus mit GIS-Objekten trainiert werden sollte, die nur einen Schlag beinhalten.

Objektklasse	Anzahl d. Objekte	enthaltene Fehler	Verwechslung zwischen Acker- und Grünland
Ackerland	328	57 (82,6%)	45 (85,9%)
Grünland	449	151 (66,4%)	33 (88,3%)
Acker-/Grünland	777	208 (73,2%)	78 (87,0%)
Andere Klassen	1699	2	-
Gesamt	2766	210	78

Tabelle 13: Statistiken über GIS-Ackerland- und Grünlandobjekte des ATKIS der Testszene HBS, bezogen auf die Anzahl der Objekte. In Klammern ist die jeweilige *TG a priori* angegeben.

Objektklasse	Fläche d. Objekte	enthaltene Fehler	Verwechslung zwischen Acker- und Grünland
Ackerland	85 km ²	2,70 km ² (96,8%)	2,63 km ² (97,0%)
Grünland ⁶	10 km ²	1,24 km ² (87,8%)	0,84 km ² (91,5%)
Acker-/Grünland	95 km ²	3,94 km ² (95,9%)	3,47 km ² (96,4%)
Andere Klassen	39 km ²	0,02 km ²	-
Gesamt	134 km²	3,96 km²	3,47 km²

Tabelle 14: Statistiken über GIS-Ackerland- und Grünlandobjekte des ATKIS der Testszene HBS, bezogen auf die Flächengröße. In Klammern ist die jeweilige *TG a priori* angegeben.

Ausgedünntes Schlagkataster

Für die Testszene Halberstadt wurde zusätzlich ein ausgedünnter Datensatz des Schlagkatasters erzeugt, der dazu benutzt wird, um speziell den zweiten Analyseschritt, also die Trennung von Acker- und Grünland näher zu untersuchen. Grundlage war der Referenzdatensatz des Schlagkatasters, wobei beim ausgedünnten Schlagkataster nur GIS-Objekte mit einer Flächengröße von mehr als 0,5 ha berücksichtigt werden. Im ausgedünnten Schlagkataster werden zunächst zufällig Trainings- und Evaluierungsobjekte gewählt (s. Abbildung 33) und den Evaluierungsobjekten anschließend zufällig künstlich Fehler hinzugefügt, um GIS-Fehler zu simulieren. Hinzufügen von zufälligen künstlichen Fehlern bedeutet, dass von zufällig gewählten GIS-Objekten die Klassenzugehörigkeit geändert wird. Eine Zusammenfassung der Statistiken dieses Datensatzes ist in Tabelle 15 gegeben. Die Angaben beziehen sich nur auf die Anzahl der Objekte, da der Fokus der Untersuchung sich darauf beschränkt.

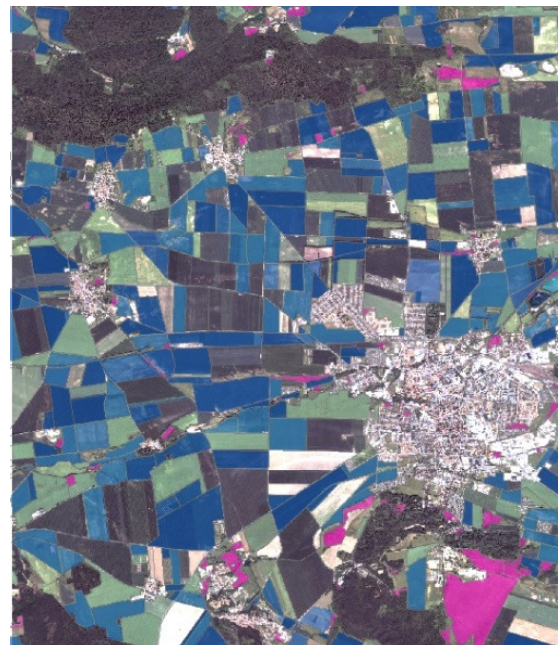


Abbildung 33: Trainingsgebiete objektbasierter Klassifikation HBS, ausgedünntes Schlagkataster (blau: Ackerland, magenta: Grünland)

Objektklasse	Anzahl d. Objekte	Verwechslung zwischen Acker- und Grünland	Trainingsgebiete für die objektbasierte Klass.
Ackerland	214	5 (97,7%)	234
Grünland	89	25 (71,9%)	70
Gesamt	303	30 (90,1%)	304

Tabelle 15: Statistiken ausgedünnter Schlagkataster der Testszene HBS zur Evaluierung des objektbasierten Analyseschrittes. In Klammern ist die jeweilige *TG a priori* angegeben.

⁶ Die Anzahl (Tabelle 13) und auch die Fläche (Tabelle 14) der GIS-Grünlandobjekte ist geringer als beim Schlagkataster (Tabelle 11 und Tabelle 12), da im Vergleich zum Schlagkataster zahlreiche Grünlandflächen nicht als einzelne GIS-Objekte erfasst wurden, sondern Bestandteil eines GIS-Ackerlandobjektes sind.

5.2.3.2. Hildesheim, Niedersachsen

Die Szene Hildesheim (kurz: HI) ist vom Stadtgebiet Hildesheim geprägt, wobei auch größere Acker- und Waldgebiete enthalten sind. Da in diesem ATKIS-Datensatz so gut wie keine Fehler vorhanden waren, wurden Fehler simuliert. Dabei wurde die Objektart von GIS-Objekten von Neubaugebieten am Stadtrand geändert (Siedlung in Grünland und vereinzelt auch in Ackerland, abhängig von den benachbarten Flächen), um ein mögliches Szenario für ein nicht aktuelles GIS zu simulieren. Statistiken zum Datensatz sind in Tabelle 16 und Tabelle 17 zusammengefasst. Die Angaben in Tabelle 16 beziehen sich wieder zunächst auf die Anzahl der GIS-Objekte, die in Tabelle 17 auf die Flächengröße. Die Bedeutung der Tabellenspalten ist identisch mit denen der Testszene Halberstadt. Es wurden nur wenige Fehler simuliert, so dass die *TG* im Allgemeinen sehr hoch ist und oft auch bereits die 95% erreicht (sowohl bezogen auf die Anzahl der GIS-Objekte als auch auf die Flächengröße). Mit Hilfe dieses Datensatzes kann das Verhalten des Verfahrens bei einer hohen *TG a priori* geprüft werden.

Objektklasse	Anzahl d. Objekte	enthaltene Fehler	Verwechslung zwischen Acker- u. Grünland	Trainingsgebiete für d. objektbasierte Klass.
Ackerland	473	29 (93,9%)	14 (96,9%)	93
Grünland	231	37 (84,0%)	15 (91,5%)	44
Acker-/Grünland	704	66 (90,6%)	29 (95,4%)	137
Andere Klassen	2722	5	-	
Gesamt	3426	71	29	

Tabelle 16: Statistiken über GIS-Ackerland- und Grünlandobjekte des ATKIS der Testszene HI, bezogen auf die Anzahl der Objekte. In Klammern ist die jeweilige *TG a priori* angegeben.

Objektklasse	Fläche d. Objekte	enthaltene Fehler	Verwechslung zwischen Acker- u. Grünland	Trainingsgebiete für objektbasierte Klass.
Ackerland	54,65 km ²	0,49 km ² (99,1%)	0,38 km ² (98,6%)	11,10 km ²
Grünland	7,91 km ²	1,60 km ² (79,8%)	0,48 km ² (92,5%)	2,77 km ²
Acker-/Grünland	62,56 km ²	2,09 km ² (96,7%)	0,86 km ² (98,6%)	18,87 km ²
Andere Klassen	55,25 km ²	0,04 km ²	-	
Gesamt	117,82 km²	2,13 km²	0,86 km²	

Tabelle 17: Statistiken über GIS-Ackerland- und Grünlandobjekte des ATKIS der Testszene HI, bezogen auf die Flächengröße. In Klammern ist die jeweilige *TG a priori* angegeben.

Für die objektbasierte Klassifikation fand das Training nur auf GIS-Objekten statt, die nur einen Schlag beinhalten. Sowohl für die objektbasierte als auch für die pixelbasierte Klassifikation wurden die Trainingsdaten daher vom Benutzer ausgewählt. Die verwendeten Trainingsdaten sind in den Abbildung 34 und Abbildung 35 dargestellt.

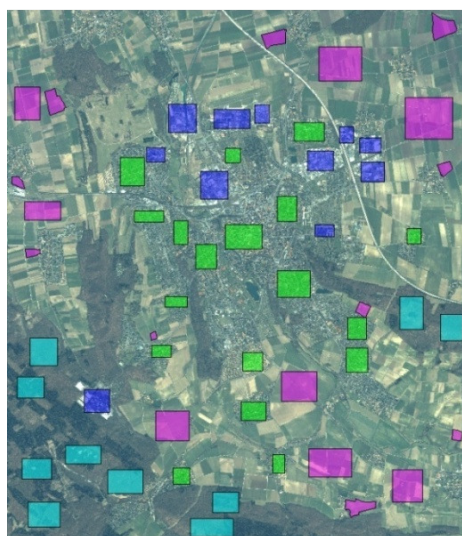


Abbildung 34: Trainingsgebiete pixelbasierte Klassifikation HI (grün: Siedlung, blau: Industrie, magenta: Wald, violett: Acker-/Grünland).



Abbildung 35: Trainingsgebiete objektbasierte Klassifikation HI (blau: Ackerland, magenta: Grünland).

5.2.3.3. Weiterstadt, Hessen

Die Szene Weiterstadt (kurz: WS) ist vor allem von landwirtschaftlichen Flächen (besonders Ackerland) geprägt; ca. 60% der Fläche ist von Ackerland bedeckt. Die Schläge innerhalb von GIS-Ackerlandobjekten in diesem Gebiet sind extrem klein, so dass mit diesem Datensatz die Grenzen der in Kapitel 4 vorgestellten Methode getestet werden kann. Auch in diesem Gebiet liegt ATKIS vor. Nähere Angaben über das GIS sind analog zur Szene Halberstadt und Hildesheim in Tabelle 18 und Tabelle 19 zusammengefasst.

Objektklasse	Anzahl d. Objekte	enthaltene Fehler	Verwechslung zwischen Acker- u. Grünland	Trainingsgebiete objektbasierte Klass.
Ackerland	686	30 (95,6%)	20 (96,9%)	189
Grünland	199	95 (52,3%)	67 (55,6%)	59
Acker-/Grünland	885	125 (85,9%)	87 (89,1%)	248
Andere Klassen	1301	27	-	
Gesamt	2186	152	29	

Tabelle 18: Statistiken über GIS-Ackerland- und Grünlandobjekte des ATKIS der Testszene WS, bezogen auf die Anzahl der Objekte. In Klammern ist die jeweilige *TG a priori* angegeben.

Objektklasse	Fläche d. Objekte	enthaltene Fehler	Verwechslung zwischen Acker- u. Grünland	Trainingsgebiete objektbasierte Klass.
Ackerland	37,58 km ²	0,51 km ² (98,6%)	0,37 km ² (99,0%)	13,96 km ²
Grünland	4,53 km ²	2,43 km ² (46,5%)	2,17 km ² (47,8%)	1,29 km ²
Acker-/Grünland	42,11 km ²	2,94 km ² (93,0%)	2,54 km ² (93,9%)	15,26 km ²
Andere Klassen	24,94 km ²	0,28 km ²	-	
Gesamt	67,05 km²	3,22 km²	2,54 km²	

Tabelle 19: Statistiken über GIS-Ackerland- und Grünlandobjekte des ATKIS der Testszene WS, bezogen auf die Flächengröße. In Klammern ist die jeweilige *TG a priori* angegeben.

Ähnlich wie bei den zuvor betrachteten ATKIS-Datensätzen ist die *TG a priori* von Ackerland höher als die von Grünland und die *TG a priori* beider Klassen zusammen liegt, sowohl bezogen auf die Anzahl der Objekte, als auch bezogen auf die Fläche, unter 95%. Auch im Datensatz Weiterstadt ist mindestens für die Klasse Ackerland die Tendenz zu erkennen, dass die *TG a priori*, bezogen auf die Flächengröße, größer als die *TG a priori*, bezogen auf die Anzahl der Objekte, ist. Für die Klasse Grünland gilt das in dem Fall nicht. Für diese Objektklasse liegt eine Mischung aus flächenmäßig kleinen und großen falschen GIS-Objekten vor.

Im Gegensatz zu den Szenen Halberstadt und Hildesheim gibt es in der Szene Weiterstadt zum einen nicht genügend GIS-Objekte, die nur aus einem Schlag bestehen, und zum anderen nicht genügend Ackerlandsegmente, die groß genug sind, um den Algorithmus der objektbasierten Klassifikation zur Trennung von Acker- und Grünland zu trainieren. Das Training in dieser Szene fand daher teilweise auch auf GIS-Objekten mit mehreren Schlägen statt. In den Experimenten soll geprüft werden, ob die in Kapitel 4 vorgestellte Methode auch unter diesen Bedingungen zufriedenstellende Ergebnisse für die Verifikation erreichen kann. Die Trainingsgebiete wurden vom Benutzer festgelegt. Die verwendeten Trainingsdaten sind in Abbildung 36 und Abbildung 37 dargestellt.

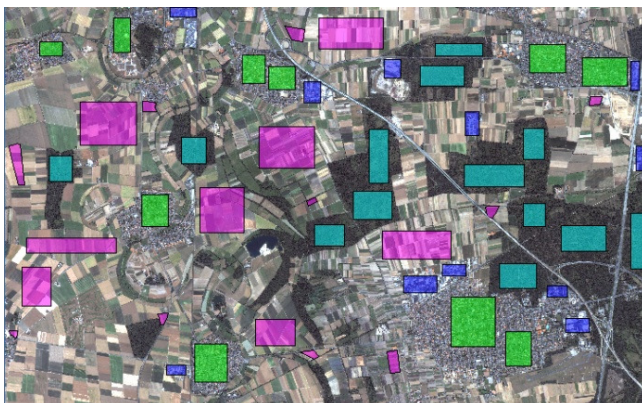


Abbildung 36: Trainingsgebiete pixelbasierte Klassifikation WS (grün: Siedlung, blau: Industrie, türkis: Wald, magenta: Acker-/Grünland).

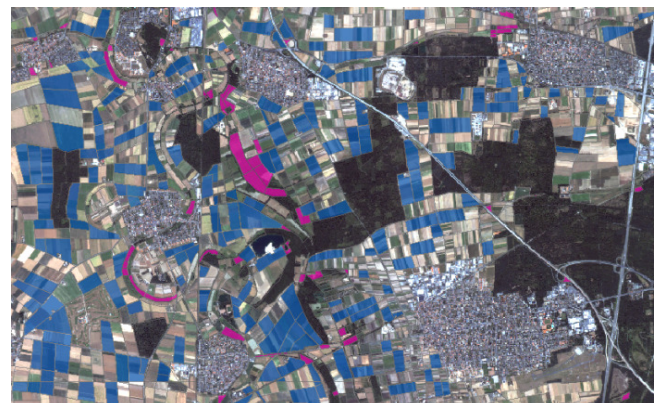


Abbildung 37: Trainingsgebiete objektbasierte Klassifikation WS (blau: Ackerland, magenta: Grünland).

5.2.3.4. Rukieten bei Rostock, Mecklenburg-Vorpommern

Im Gegensatz zu allen anderen Testgebieten handelt es sich bei dem Testgebiet Rukieten bei Rostock (kurz: RUK) um ein sehr kleines Gebiet, welches nur ca. 18 km² groß ist. Dafür liegen für dieses Testgebiet multitemporale Bilddaten vor, die zu 12 verschiedenen Zeitpunkten innerhalb von zwei Jahren aufgenommen wurden. Die Bilddaten wurden mit dem Niedrigpreissystem PFIFF (Grenzdörffer, 2005) erfasst. Für die Evaluierung wurden nur Zeitpunkte verwendet, an denen vier Bildkanäle zur Verfügung standen (s. Tabelle 20).

Flug	Aufnahmezeitpunkt	Kanäle
2	07.10.2005	NIR, R, G, B
3	23.11.2005	NIR, R, G, B
4	17.04.2006	NIR, R, G, B
5	04.05.2006	NIR, R, G, B
6	09.06.2006	NIR, R, G, B
7	25.06.2006	NIR, R, G, B

Tabelle 20: Bilddaten für das Testgebiet Rukieten bei Rostock.

Im vorhandenen Schlagkataster dieses Gebietes sind nur GIS-Objekte der Objektarten Acker- und Grünland enthalten. Für die Untersuchungen wurden nur GIS-Objekte verwendet, deren Objektart sich innerhalb des Untersuchungszeitraumes nicht geändert hat. Eine Übersicht über die enthaltenen GIS-Objekte ist in Tabelle 21 gegeben. Ungefähr die Hälfte aller Objekte wurde zufällig als Trainingsobjekte, die andere als Evaluierungsobjekte festgelegt. Die Trainingsgebiete für die objektbasierte Klassifikation sind graphisch in Abbildung 38 dargestellt. Da das Schlagkataster fehlerfrei ist, wurden in den Evaluierungsobjekten zufällig vier GIS-Ackerlandobjekte als GIS-Grünlandobjekte und sechs GIS-Grünlandobjekte als GIS-Ackerlandobjekte gekennzeichnet. Dies entspricht einer Fehlerquote von etwas mehr als 10%. Wie Tabelle 21 zu entnehmen ist, liegt die *TG a priori* somit immer unter der geforderten Marke von 95%. Wie bei dem ausgedünnten Schlagkataster der Testszene Halberstadt wird auch hier auf eine Darstellung der Flächenanteile verzichtet, da auch dieser Datensatz nur für eine Spezialuntersuchung des objektbasierten Analyseschrittes verwendet wird.

Obwohl nur eine sehr geringe Anzahl an GIS-Objekten für diesen Datensatz zur Verfügung stehen und es wie bereits angesprochen Kalibrierungsproblemen des NIR-Kanals gibt, soll dieser Datensatz genutzt werden, um die in Kapitel 4 eingeführte Methode zu evaluieren und eventuelle Grenzen aufzuzeigen.



Abbildung 38: für die objektbasierte Klassifikation RUK (blau: Ackerland, magenta: Grünland).

Objektklasse	Anzahl d. Objekte	Verwechslung zwischen Acker- und Grünland	Trainingsgebiete für die objektbasierte Klass.
Ackerland	21	4 (81,0%)	24
Grünland	17	6 (64,7%)	14
Gesamt	38	10 (73,7%)	38

Tabelle 21: Statistiken über GIS-Ackerland- und Grünlandobjekte des ATKIS der Testszene RUK, bezogen auf die Anzahl der Objekte. In Klammern ist die jeweilige *TG a priori* angegeben.

5.3. Wahl der Parameter

Im vorliegenden Abschnitt sollen alle Parameter, die für den in Kapitel 4 vorgestellten Ansatz benötigt werden, aufgeführt und deren Festlegungen erörtert werden. Es wird gezeigt, dass viele Parameter automatisch gelernt werden können oder durch den jeweiligen GIS-Objektartenkatalog vorgegeben sind. Nur wenige Parameter sind auf Grund von Erfahrungswerten durch den Benutzer einzustellen.

5.3.1. Parameter für die Klassifikation zur Identifizierung von GIS-Ackerland- und Grünlandobjekten

Parameter, die für den ersten Analyseschritt, der pixelbasierten Klassifikation mit Überführung des Klassifikationsergebnisses auf GIS-Objekte (Abschnitt 4.3.), notwendig sind, hängen ab von:

- den zu unterscheidenden Klassen,
- dem verwendeten GIS und
- der geometrischen Auflösung der Bilddaten.

In dieser Arbeit werden die Klassen *Acker-/ Grünland*, *Wald*, *Siedlung* und *Industrie* auf 1 m pansharpend IKONOS-Bilder für die Verifikation von ATKIS bzw. eines Schlagkatasters klassifiziert. Die Luftbilder werden nicht verwendet, da im dazu vorliegenden GIS nur die Klasse *Acker-* und *Grünland*, aber nicht die Klassen *Wald*, *Siedlung* und *Industrie* vorhanden sind und somit die gemeinsame Klasse *Acker-/ Grünland* nicht von anderen Klassen getrennt werden muss; die Durchführung des ersten Analyseschrittes ist damit nicht nötig. Die Parameter für die Klassifikation und die anschließende Übertragung des pixelbasierten Klassifikationsergebnisses auf ein GIS-Objekt sind in Tabelle 22 zusammengefasst.

Parameter 1, also die Pyramidenstufe der Bildpyramide, auf der die pixelbasierte Klassifikation mittels des Ansatzes von Gimel'farb (1996) durchgeführt wird, wird mit Hilfe des Ansatzes von Becker et al. (2009) automatisch für jede zu klassifizierende Objektklasse separat trainiert. Die erste Pyramidenstufe entspricht der Auflösung des Originalbildes. Die pixelbasierte Klassifikation wurde auf den IKONOS-Bildern der Testszenen Halberstadt, Hildesheim und Weiterstadt durchgeführt. Dabei wurde fast immer die Pyramidenstufe 4 (entspricht einer geometrischen Auflösung von 8 m) genutzt. Die Ausnahmen sind die Pyramidenstufen für die Klassifikation der Klassen *Siedlung* und *Wald* für das IKONOS-Bild der Testszene Halberstadt. Hier liegt die Pyramidenstufe 3 vor (entspricht einer geometrischen Auflösung von 4 m).

Nr.	Parameter	Bestimmung	Wert
1	Geeignete Auflösungsstufe der Bildpyramide für die pixelbasierte Klassifikation nach (Gimel'farb, 1996)	Voll automatisch mittels Training	meist 4, selten 3
2	Schwellwert s_q für den Quotient aus Gleichung 4.5	Benutzerdefiniert	0,3 oder 0,4 (s. 5.4)
3	Schwellwert s_A für die Mindestgröße eines kompakten Fehlers	mittels Objektartenkatalog	0,1ha (ATKIS)
4	Schwellwert s_E für die Anzahl der Schritte eines morphologischen Filters (Erosion) E , die benötigt werden, bis die Region eines vermeintlichen kompakten Fehlers verschwindet	Benutzerdefiniert	20

Tabelle 22: Parameter für den ersten Analyseschritt nach Abschnitt 4.3.

Mit dem Parameter 2, dem Schwellwert s_q , kann beeinflusst werden, wie sensibel das System eingestellt werden soll. Ist dieser Parameter sehr klein, ist nur ein sehr geringer Anteil an falschen Pixeln in einem GIS-Objekt erlaubt. Mit *falschen Pixeln* werden Pixel bezeichnet, die nicht der im GIS vorgegebenen Klasse übereinstimmen. Die Wahl des Parameters hängt davon ab, wie häufig in einem GIS-Objekt einer bestimmten Klasse auch Objekte anderer Klassen zu finden sind. In einem Siedlungsobjekt ist es zum Beispiel nicht selten, dass auch Bäume oder Grünflächen zu erkennen sind, die laut GIS-Objektartenkatalog in einer Siedlungsfläche auch erlaubt sind. Der Wert für s_q ist somit für diese Klasse relativ hoch zu wählen. Für Acker- und Grünland hingegen ist der Wert für s_q niedriger anzusetzen. Erfahrungswerte sind 30% bis 40%, denn auch hier sind z.B. Bäume als Ackerlandbegrenzung oder versiegelte Flächen bei straßenbegleitendem Grünland möglich. Eine detaillierte Analyse des Schwellwertes und des Einflusses auf das Verifikationsergebnis wird im ersten Experiment in Abschnitt 5.4 beschrieben.

Ein GIS-Objekt wird aber nicht nur dann abgelehnt, wenn der Schwellwert s_q überschritten wird, sondern auch, wenn ein kompakter Fehler vorliegt. Ein kompakter Fehler liegt vor, wenn die Schwellwerte für Parameter 3 und Parameter 4 gleichzeitig überschritten werden. Parameter 3 legt die Anzahl der Pixel fest, die mindestens nötig sind, damit ein kompakter Fehler vorliegt. Dieser Wert kann direkt an Hand des Objektartenkatalogs festgelegt werden und richtet sich nach den Mindestkartierflächen im Objektartenkatalog der Objektarten Siedlung, Industrie und Wald (Tabelle 8). Der Schwellwert wird daher auf 1000 m² (0,1 ha) festgelegt. Parameter 4, der Schwellwert s_E für die Anzahl der Schritte eines morphologischen Filters (Erosion) E , die benötigt werden, bis die Region eines vermeintlichen kompakten Fehlers verschwindet, hängt vom Objektartenkatalog und Erfahrungswerten ab. Wenn z.B. die

Mindestkartierfläche 1 ha beträgt, ist es unwahrscheinlich, dass ein kompakter Fehler schmäler als 40 m ist (Erfahrungswert). Der Schwellwert s_E wurde daher auf Grund der geometrischen Auflösung von 1 m auf 20 festgesetzt. Ein kompakter Fehler liegt also nur dann vor, wenn die Region von verbundenen falschen Pixeln, die ein und derselben Klasse angehören, mindestens 40 m breit ist und der kompakte Fehler die Mindestfläche (Parameters 3) überschreitet.

5.3.2. Parameter für die Segmentierung

Die nächste Gruppe von Parametern sind die der Segmentierung (Abschnitt 4.4). Ziel der Segmentierung ist das Finden homogener Segmente innerhalb eines GIS-Acker-/Grünlandobjektes, die unterschiedliche Schläge repräsentieren. Die Parameter der Segmentierung hängen von dem Bilddaten ab und die des Nachbearbeitungsschrittes der Segmentierung vom GIS. Während der Experimente wurden 1 m pansharping IKONOS-Bilder benutzt; als GIS lag ATKIS vor. Die Parameter für die Segmentierung sind in Tabelle 23 zusammengefasst.

Nr.	Parameter	Bestimmung	Wert
5	Parameter σ der Gaußfunktion, der den Grad der Glättung vor dem Anwenden des Wasserscheidenverfahrens bestimmt	Abhängig von den Bilddaten	s. Tabelle 24
6	Signifikanzniveau α		s. Tabelle 24
7	Schwellwert s_T für den maximalen Wert von T_{ij} mit $s_T \in [0\%, 100\%]$ als Stärke der Kante zwischen zwei Segmenten		s. Tabelle 24
8	Parameter k_D für die Bestimmung des Schwellwerts s_D nach Gleichung (4.12) für den max. Wert von D_{ij} aus Gleichung (4.7)		s. Tabelle 24
9	Parameter k_F für die Bestimmung des Schwellwerts s_F nach Formel (4.13) für den maximalen Wert von F_{ij} aus Gleichung (4.10)		s. Tabelle 24
10	Nachbearbeitung (Verschmelzen kleiner Regionen mit der Umgebung)	Objektarten-katalog	1 ha (ATKIS)

Tabelle 23: Parameter für die Segmentierung nach Abschnitt 4.4.

Die Parameter 5 bis 9 hängen nicht nur von der geometrischen Auflösung der Bilddaten ab, sondern auch von den spektralen Merkmalen der in der Szene enthaltenen verschiedenen Ackerland- und Grünlandschläge. Experimente haben gezeigt, dass die Parameter auf alle GIS-Objekte innerhalb einer IKONOS-Szene gleichermaßen angewendet werden können.

Die Parameter 5 bis 9 wurden zunächst empirisch für einen kleinen Ausschnitt der Testszene Halberstadt ermittelt. Dazu waren Tests notwendig. Die Beurteilung erfolgte mittels visueller Kontrolle. Für die Testszene Hildesheim werden die Parameter unverändert übernommen, während die Parameter für die Testszene Weiterstadt geringfügig geändert werden mussten (Parameter k_F ist strenger eingestellt). Die Parameter 5-9 sind detailliert für jede Testszene in Tabelle 24 zusammengefasst.

Parameter	Halberstadt	Hildesheim	Weiterstadt
σ	1	1	1
α	10%	10%	10%
s_T	80%	80%	80%
k_D	1	1	1
k_F	1,25	1,25	1,5

Tabelle 24: Verwendete Parameter der Segmentierung für die einzelnen Testszenen.

Der Wert des Parameters 10 wird durch den GIS-Objektartenkatalog definiert. Wenn nach der Segmentierung ein Segment vorliegt, das nur einen weiteren Nachbarn hat und kleiner als die Mindestkartierfläche von Acker-/Grünland ist, so wird dieses Segment der angrenzenden Fläche zugeordnet. Da die Segmentierung Bestandteil des zweiten Analyseschrittes ist, also der Trennung zwischen Acker- und Grünland, kann davon ausgegangen werden, dass es sich bei dem mit den Nachbarn zu verschmelzenden Segment um ein Ackerland- oder Grünlandsegment handelt. Deshalb wird die Mindestkartierfläche von Acker- und Grünland verwendet (1 ha bei ATKIS).

5.3.3. Parameter für die Klassifikation eines Segmentes

Die Parameter, die für die Klassifikation eines Segmentes notwendig sind (Abschnitt 4.5), gliedern sich in zwei Gruppen. In der ersten Gruppe sind die Parameter für die Bestimmung der Merkmale enthalten. Bei der

zweiten Gruppe handelt es sich um Parameter für die Klassifikation mittels SVM. Die Parameter aus beiden Gruppen werden getrennt betrachtet.

Die Parameter der ersten Gruppe (Abschnitt 4.5.1) sind identisch für alle in den Experimenten verwendeten Testszenarien und sind in Tabelle 25 zusammengefasst. Der Parameter 11 ist ein Erfahrungswert dafür, um wie viel ein Segment vor der Analyse reduziert werden muss, damit Unregelmäßigkeiten am Segmentrand die Klassifikation nicht beeinflussen. Dieser Parameter hängt vom zu untersuchenden Gebiet und den dort vorkommenden landwirtschaftlichen Praktiken ab. Er wurde einheitlich für alle Tests auf 15m gesetzt. Der Parameter 12 ist ein Erfahrungswert dafür, ab welcher Flächengröße eine zuverlässige Bestimmung der Parameter möglich ist. Dies gilt insbesondere für die strukturellen Merkmale. Die Parameter 13 und 14, die für die Bestimmung der textuellen Merkmale notwendig sind, sind auf in der Literatur übliche Werte gesetzt (Tönnies, 2005). Parameter 15 wurde über empirische Tests bestimmt.

Nr.	Parameter	Bestimmung	Wert
11	Bereich, um die ein Segment reduziert wird (Erosion)	Erfahrungswert	15 m
12	Mindestsegmentgröße für d. zuverlässige Berechnung d. Merkmale	Erfahrungswert	0,5 ha
13	Aus der Gruppe der textuellen Merkmale: Δ aus (Gleichung 4.19)	Literatur	1
14	Aus der Gruppe der textuellen Merkmale: α aus (Gleichung 4.19)	Literatur	0°, 45°, 90°, 135°
15	Aus der Gruppe der strukturellen Merkmale: unterer und oberer Hystereseschwellwert	Empirische Tests	50/70

Tabelle 25: Parameter für die Bestimmung der Merkmale für die objektbasierte Klassifizierung eines Segmentes nach Abschnitt 4.5.1.

Die Parameter für die Klassifikation mittels der SVM (Abschnitt 4.5.2) sind in Tabelle 26 zusammengefasst. Bereits im Abschnitt 3.1.1 wurde herausgearbeitet, warum die RBF-Kernfunktion gewählt wurde. Alle weiteren Parameter für die Klassifikation können automatisch mittels Gittersuche unter Verwendung der Kreuzvalidierung, wie in Abschnitt 3.1.4 beschrieben, trainiert werden. Auf eine Auflistung aller Parameter für die Tests wird an dieser Stelle verzichtet, da es sich um insgesamt 16 verschiedene Testszenarien handelt und die Werte ohnehin automatisch bestimmt werden.

Nr.	Parameter	Bestimmung	Wert
16	Wahl der Kernelfunktion	Stand der Wissenschaft	RBF
17	Parameter C (Strafwert für Fehlklassifikationen, s. Gleichung (3.25))	Automatisch mittels Gittersuche und Kreuzvalidierung	
18	Parameter γ der RBF-Kernfunktion(s. Gleichung (3.20))		

Tabelle 26: Parameter für die Klassifikation der Merkmale für die objektbasierte Klassifizierung eines Segmentes nach Abschnitt 4.5.2.

5.3.4. Parameter für die Klassifikation und Verifikation eines GIS-Objektes

Im letzten Schritt für die Klassifikation und Verifikation eines GIS-Objektes wird die Summe aller Segmentflächen, die entweder der Klasse Acker- oder Grünland zugewiesen wurden, betrachtet (Abschnitt 4.6.). Das GIS-Objekt wird der Objektklasse zugeordnet, die die Mindestkartierfläche überschritten hat. Werden von beiden oder von keiner Objektklasse die Mindestkartierfläche erreicht, wird das Objekt einer Zurückweisungsklasse zugeordnet. Der Wert der Mindestkartierfläche ergibt sich aus dem vorliegenden GIS. Bei ATKIS beträgt dieser 1 ha (s. Tabelle 27). Für diesen Analyseschritt sind keine weiteren Parameter nötig.

Nr.	Parameter	Bestimmung	Wert
19	Parameter über Annahme/Ablehnung eines GIS-Acker-/Grünlandobjektes durch das System	Objektartenkatalog	1 ha (ATKIS)

Tabelle 27: Parameter für die Klassifikation der Merkmale für die objektbasierte Klassifizierung eines Segmentes nach Abschnitt 4.6.

5.4. Evaluation des ersten Analyseschrittes

Im ersten Versuch soll überprüft werden, ob der im Abschnitt 4.3 vorgestellte Algorithmus für die Trennung der gemeinsamen Klasse Acker-/Grünland gegenüber anderen Objektarten für die Aufgabe der Verifikation geeignet ist. Im Folgenden werden nur die Verifikationsergebnisse präsentiert, bei denen das pixelbasierte Klassifikationsergebnis bereits auf die GIS-Objekte übertragen wurde. Auf eine explizite Evaluierung des

pixelbasierten Klassifikationsergebnisses wird verzichtet, da der Fokus auf der Verifikation von GIS-Objekten liegt. Implizit wird der Erfolg der pixelbasierten Klassifikation jedoch bei der Bewertung des Verifikationsergebnisses der GIS-Objekte ebenfalls überprüft.

Wie in 5.3.1 beschrieben, ist vor allem eine genauere Untersuchung des Parameters 2 (Schwellwert s_q für den Quotient aus Gleichung (4.5)) von Interesse. Die Untersuchungen zu der Eignung des in Abschnitt 4.3 beschriebenen Verfahrens erfolgt in zwei Schritte. Zunächst wird mittels eines Testgebietes analysiert, welche Werte für s_q geeignet sind. Diese werden dann im zweiten Schritt auf weitere Gebiete übertragen und es wird getestet, ob diese Schwellwerte auch auf anderen Szenen ohne Anpassung angewendet werden können.

5.4.1. Untersuchung der Abhängigkeit des Verifikationsergebnisses von s_q

Wie erläutert, wird im ersten Experiment der Erfolg der Verifikation in Abhängigkeit von s_q für eine Testszene näher untersucht. Für diesen Test wird eine Testszene benötigt, bei der neben GIS-Ackerland-/Grünlandobjekten auch GIS-Objekte anderer Klassen wie Siedlung, Industrie und Wald vorkommen. Die Testszenen Rukieten sowie Halberstadt mit ausgedünntem Schlagkataster sind daher ungeeignet. Außerdem sollte die geforderte Genauigkeit von 95% *TG a priori* nicht erfüllt sein, da sonst keine zuverlässige Aussage darüber gemacht werden kann, ob und wann mit Hilfe des pixelbasierten Verfahrens die gewünschte *TG a posteriori* erzielt werden kann. Für dieses Experiment wird daher die Testszene Halberstadt unter Verwendung von ATKIS benutzt. In der Testszene Halberstadt liegt die *TG a priori* der gemeinsamen Klasse Acker-/Grünland bei 82,2%, bezogen auf die Anzahl der Objekte. Bezogen auf die Fläche, liegt die *TG a priori* dieser Szene bei über 95%. Da nur die *TG a priori*, bezogen auf die Anzahl der Objekte kleiner als 95% ist, wird die Darstellung der Ergebnisse dieses Experimentes auf die Auswertung auf Basis der Anzahl der Objekte beschränkt. Für die Testreihe werden die Schwellwerte für s_q von 0% bis 100% in 10%-Schritte variiert und für jeden Schwellwert die *TG a posteriori* sowie die *Effizienz*, der relative Anteil der *False Positives*, die *error detection rate* und die *overall accuracy* bestimmt. Die Ergebnisse sind in Abbildung 39 zusammengefasst.

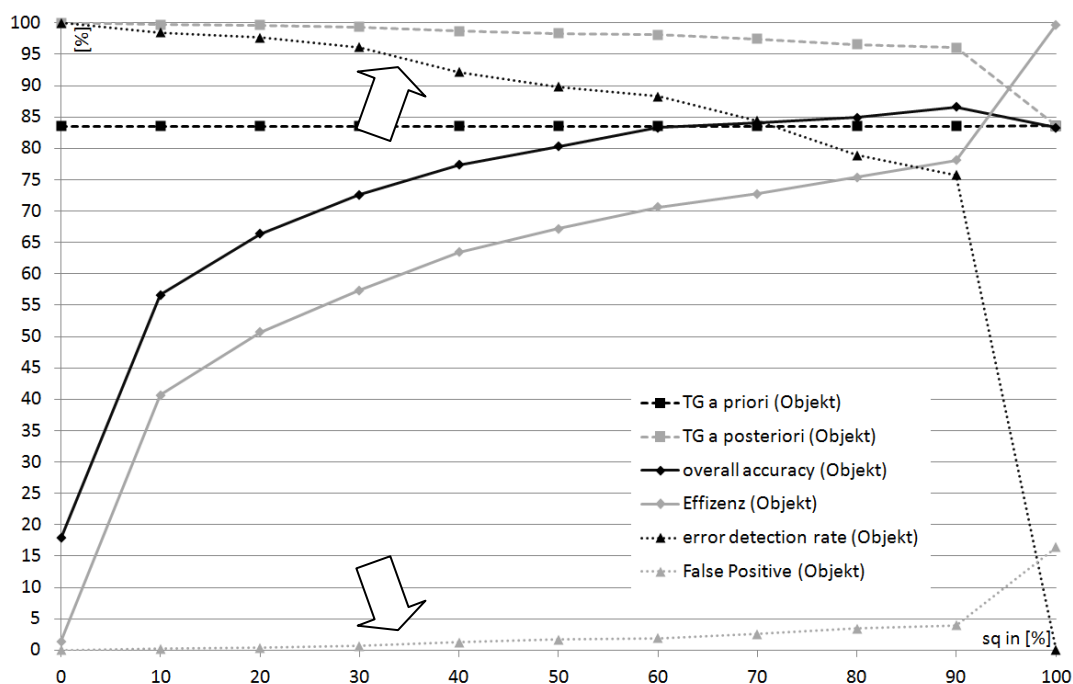


Abbildung 39: Untersuchung der Anhängigkeit des Verifikationsergebnisses von s_q für die Verifikation der gemeinsamen Klasse Acker-/Grünland des ersten Analyseschrittes für die Testszene HBS (IKONOS, ATKIS).

Beträgt der Wert für $s_q = 0\%$, sind keinerlei falsche Pixel in einem GIS-Objekt erlaubt. Dieser Fall trifft auf nur 11 sehr kleine GIS-Objekte (1,4%) zu. Ein Wert von $s_q = 0\%$ bedeutet, dass alle Fehler im GIS gefunden werden (*error detection rate* = 100%), weil so gut wie alle GIS-Objekte vom System abgewiesen werden und daher so gut wie jedes GIS-Objekt manuell nachbearbeitet werden muss. Die *Effizienz* ist daher sehr niedrig.

Erhöht sich der Wert für s_q , sind mehr falsche Pixel zugelassen. Während die Anzahl der Fehler, die im GIS verbleiben (*False Positive*), und die *Effizienz* steigen, sinkt die *error detection rate* zeitgleich. Obwohl die Anzahl der *False Positive* steigt, unterschreitet die *TG a posteriori* den gewünschten Wert von 95% selbst dann nicht, wenn $s_q = 90\%$ beträgt. Bei einem Wert von $s_q = 100\%$ werden alle GIS-Objekte, die keinen kompakten Fehler beinhalten, vom System angenommen. Bei der hier verwendeten Testszene trifft dies nur bei zwei GIS-Objekten (0,3%) nicht zu. Die *Effizienz* ist dann sehr hoch, während die *TG a posteriori* gegenüber der *TG a priori* gleich ist (*error detection rate* = 0%) und daher keine Verbesserung der *TG* erreicht werden kann. Während die *overall accuracy* genauso wie die *Effizienz* bei einem größer werdenden s_q steigt, sinkt sie bei $s_q = 100\%$ wieder ab, da die *overall accuracy* die Summe aller korrekt verifizierten Objekte ist und der Anteil der korrekt erkannten Fehler im GIS auf Null sinkt.

Bei der Betrachtung der Ergebnisse des ersten Analyseschrittes ist zu beachten, dass der gesamte Algorithmus niemals besser als der erste Analyseschritt werden kann. Eine *TG* von 95% mit dem ersten Analyseschritt zu erreichen, ist daher nicht ausreichend. Das Ergebnis muss besser sein. Ein *TG*, wie sie mit $s_q = 50\%$ erreicht wird (*TG* = 98,3%), ist nach Erfahrungswerten noch brauchbar. Während eine Fehlerquote von 50% erfahrungsgemäß für Grünland realistisch ist (viele andere Objekte wie Versiegelungsflächen und Bäume sind gerade bei Verkehrsbegleitgrün keine Seltenheit), ist eine Fehlerrate für die gemeinsame Klasse Acker-/Grünland mit 50% zu hoch, 30% oder 40% sind realistischer.

In Abbildung 39 ist zwischen $s_q = 30\%$ und $s_q = 40\%$ bezüglich der *False Positive*, *TG a posteriori* und *error detection rate* ein Sprung zu erkennen (s. Pfeile), der deutlicher als bei anderen Intervallen ausfällt. Es ist interessant zu untersuchen, ob ein ähnliches Verhalten auch in den anderen Testszenen zu beobachten ist. Aus diesem Grund und weil erfahrungsgemäß ein Wert für s_q von 30% bzw. 40% realistisch ist, werden diese beiden ausgewählten Werte auch auf den anderen Szenen getestet.

5.4.2. Untersuchung des Erfolgs der Verifikation für zwei verschiedene s_q in unterschiedlichen Testszenen

Im zweiten Teil des ersten Experiments soll nun untersucht werden, ob die Verifikationsergebnisse für die zuvor bestimmten Werte für $s_q = 30\%$ und $s_q = 40\%$ ebenfalls für andere Testszenen zufriedenstellende Ergebnisse liefern. Für diese Testreihe werden alle Testszenen genutzt, für die ein GIS mit Ackerland- und Grünlandobjekten, aber auch Siedlungs-, Industrie- und Waldobjekten vorliegt und die Anzahl der GIS-Objekte groß ist. Dies trifft auf die Testszenen Halberstadt (mit Schlagkataster und ATKIS) sowie Hildesheim (ATKIS) und Weiterstadt (ATKIS) zu. Die Ergebnisse sind in Abbildung 40 zu sehen.

Zunächst ist festzustellen, dass in allen Testszenarien eine *TG a posteriori* von über 95% erreicht werden konnte, wobei diese in der Testszene Weiterstadt a priori bereits vorlag. Dass die *TG a priori* von 95% für alle Testszenarien unter Betrachtung der Fläche bereits erreicht ist, unter Betrachtung der Anzahl der Objekte jedoch nicht, spricht ebenfalls dafür, dass es sich bei den fehlerhaften GIS-Objekten um sehr kleine GIS-Objekte handelt. Ein weiterer Hinweis, dass die fehlerhaften Objekte besonders klein sind, erkennt man bei der Betrachtung des Unterschiedes der *Effizienz* von objektweiser Betrachtung und unter Berücksichtigung der Fläche. Die *Effizienz*, bezogen auf die Anzahl der Objekte, ist immer wesentlich geringer gegenüber der *Effizienz*, die auf der Flächengröße der GIS-Objekte basiert. Ähnliche Beobachtungen bezüglich der Flächengröße fehlerhafter GIS-Objekte werden auch in nachfolgenden Experimenten beobachtet. Die *TG a priori* für ein bestimmtes GIS muss konstant bleiben, da die Anzahl der Fehler im GIS konstant ist (verschiedene Experimente, die mit einem GIS durchgeführt werden, beeinflussen die Anzahl der Fehler im GIS natürlich nicht). Eine konstante *TG a priori* ist in Abbildung 40 für jedes GIS gut zu erkennen. Der Grund dafür, dass die *TG a posteriori* sinkt und die *Effizienz* steigt, wenn der Schwellwert s_q erhöht wird, wurde bereits im vorhergehenden Abschnitt erörtert. Beide Charakteristika sind in Abbildung 40 zu sehen.

Die *Effizienz* in der Testszene Halberstadt ist sehr viel geringer als in den anderen Testszenen und in dieser Szene besonders bei der Verwendung von ATKIS, was auf die, wie bereits erörtert, kleinen fehlerhaften GIS-Objekte zurückzuführen ist. Betrachtet man die fünf größten GIS-Objekte nicht mit, die vom System verworfen wurden, beträgt die durchschnittliche Fläche weniger als 1,5 ha.

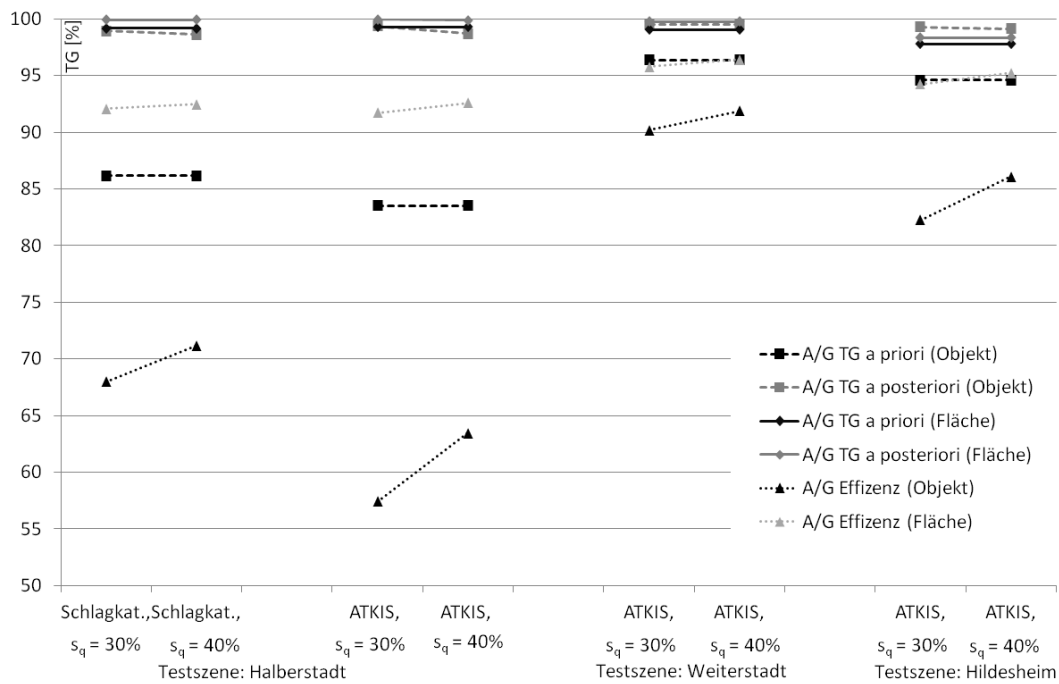


Abbildung 40: Evaluierungsergebnisse der Verifikation der gemeinsamen Klasse Acker-/Grünland (A/G) für die Testszenen Halberstadt, Weiterstadt und Hildesheim für $s_q = 30\%$ und $s_q = 40\%$.

Das schlechteste Ergebnis für die *TG a posteriori* wurde für die Testszene Halberstadt in Kombination mit dem Schlagkataster und einem Wert für s_q von 40% erreicht. Die Ergebnisse dieser Szene sind graphisch in Abbildung 41 (rechts) zu sehen. Alle dunkelgrünen (*True Positives*) und hellroten (*False Positive*) GIS-Objekte werden vom menschlichen Bearbeiter nicht mehr kontrolliert, da sie vom System akzeptiert wurden. Die hellroten GIS-Objekte sind dabei die Fehler, die im GIS verbleiben, da diese vom System nicht erkannt wurden. Sie sind in Abbildung 41 (rechts) auf Grund ihrer Größe mit Pfeilen gekennzeichnet. Alle dunkelroten Objekte sind Fehler, die im GIS gefunden werden konnten, während hellgrüne Objekte vom menschlichen Bearbeiter manuell kontrolliert werden müssen, obwohl sie richtig sind.

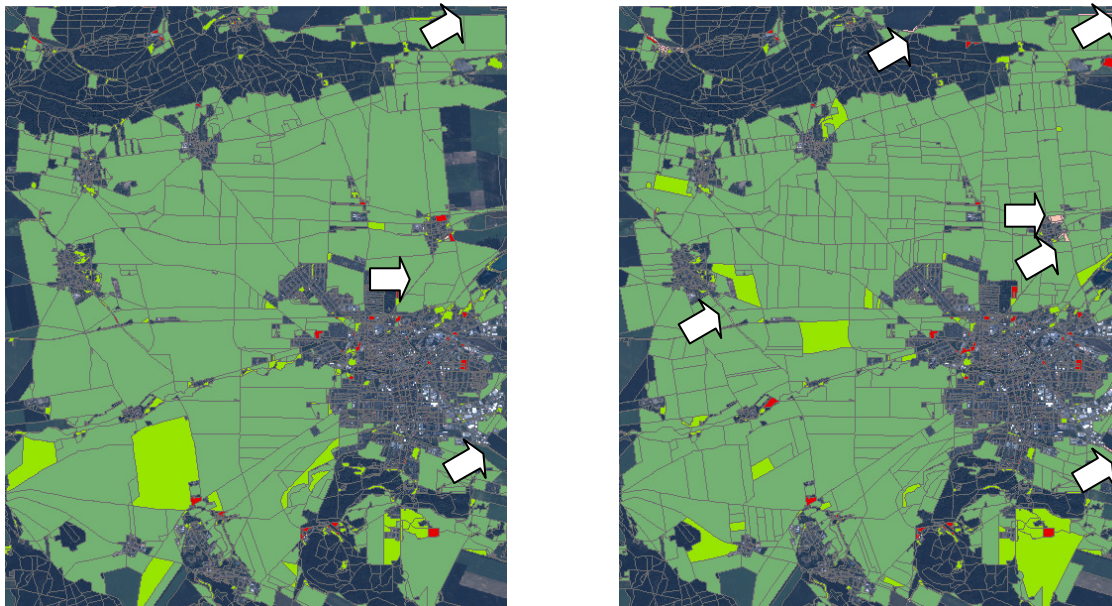


Abbildung 41: Graphische Darstellung der Konfusionsmatrix der Verifikationsergebnisse des ersten Analyseschrittes (pixelbasierte Klassifikation), links: Halberstadt, ATKIS, $s_q = 30\%$, rechts: Halberstadt, Schlagkataster, $s_q = 40\%$ (TP: dunkelgrün, FN: hellgrün, FP: hellrot (durch Pfeile gekennzeichnet), TN: dunkelrot).

Es gibt 13 *False-Positive*-Objekte (hellrot) von insgesamt 947 Ackerland-/Grünlandobjekten mit insgesamt 131 fehlerhaften GIS-Objekten der gemeinsamen Klasse Acker-/Grünland. Bei 9 der *False-Positive*-Objekte handelt es sich um sehr schmale und lange GIS-Objekte entlang von Straßen, die im GIS der Klasse Grünland zugeordnet sind, in der Referenz aber als Siedlungs- oder Industriefläche geführt werden, da der größte Teil der Fläche versiegelt ist. Zwei *False Positive* sind sehr kleine Objekte (200m^2 und 400m^2), die von Wald bedeckt sind, sich aber in der Nähe eines Ackerlandobjektes befinden. Diese beiden Objekte wurden größtenteils als Ackerlandobjekt klassifiziert und daher vom System nicht verworfen. Die letzten beiden *False-Positive*-Objekte sind GIS-Objekte, auf denen Bebauung entstanden ist, die aber zu klein ist, um ein kompakter Fehler zu sein und im Verhältnis zum Rest des GIS-Objektes ebenfalls zu klein ist, so dass q kleiner als s_q ist, wodurch diese fehlerhaften GIS-Objekte nicht detektiert werden konnten.

Auf die explizite Darstellung der *False Positives*, der *overall accuracy* und der *error detection rate* soll an dieser Stelle verzichtet werden. Die *False-Positive-Rate*, sowohl bezogen auf die Anzahl der Objekte als auch auf die Fläche, liegt fast durchweg unter 0,3% (Ausnahme ist Hildesheim mit einen Wert von weniger als 1,8%). Auch für die *overall accuracy* konnten ähnlich gute Ergebnisse erreicht werden. Sie liegt bei Betrachtung der Anzahl der Objekte zwischen 72% und 95%, sowie bei Betrachtung der Fläche zwischen 92% und 97%. Das ist ein weiterer Hinweis darauf, dass es sich bei den fehlerhaften GIS-Objekten um sehr kleine Objekte handelt. Die *error detection rate*, bezogen sowohl auf die Anzahl der Objekte als auch auf die Fläche, liegt bei über 75%. Eine Ausnahme bildet die *error detection rate* bezogen auf die Fläche der Szene Hildesheim, die mit der schlechtesten *False-Positive-Rate* für diese Konfiguration. Hauptgrund für die schlechte *error detection rate*/geringe *False-Positive-Rate* ist das GIS-Grünlandobjekt, das in Abbildung 42 zu sehen ist. In Abbildung 42 rechts ist das Klassifikationsergebnis der pixelbasierten Analyse nach Gimel'farb (1996) gezeigt. Es ist zu sehen, dass das Gebäude als Industriefläche (gelb) erfolgreich detektiert wird (s. Pfeil), das zur Ablehnung durch das System führen sollte. Da das Gebäude aber sehr klein ist und daher nicht als kompakter Fehler erkannt wird und da das Grünlandobjekt so groß ist und das Haus daher nur eine kleine Fläche im Verhältnis zu der restlichen Fläche des GIS-Objektes bedeckt, geht die Information bei der Übertragung des pixelbasierten Klassifikationsergebnisses auf das GIS-Objekte nach Busch et al. (2005) (Abschnitt 4.3.2) verloren. Eine solche Konfiguration stellt eine Ausnahme dar.



Abbildung 42: Fehlerhaftes GIS-Grünlandobjekt, das tlw. von Siedlung bedeckt ist (Testszene: HI). Links: RGB-Bild, rechts: Ergebnis d. pixelbasierten Klassifizierung (braun: Acker-/Grünland, gelb: Industrie, rot: Siedlung, grün: Wald).

5.4.3. Bewertung

Ziel dieses Analyseschrittes ist es, GIS-Ackerland- und Grünlandobjekte zu identifizieren, in denen eine Fläche oder Teilfläche anderer Nutzung vorkommt. Diese GIS-Objekte sollen vom System gekennzeichnet und dem menschlichen Bearbeiter zur Beurteilung gezeigt werden. Dabei ist es nötig, eine *TG a posteriori* von mindestens 95% zu erreichen, da der gesamte Algorithmus niemals besser als der erste Analyseschritt werden kann. Die geforderten Ziele konnten mit dem in Abschnitt 4.3 vorgestellten pixelbasierten Klassifikationsverfahren von Gimel'farb (1996) mit der anschließenden Übertragung des pixelbasierten Klassifikationsverfahrens auf GIS-Objekte nach Busch et al. (2005) erreicht werden, wie in Abbildung 40 zu sehen ist. Eine Ausnahme mit großem Einfluss auf das Ergebnis liegt in der Testszene Hildesheim vor (Abbildung 42). Diesem Problem kann man begegnen, indem die Schwellwerte sowohl für den kompakten Fehler als auch für s_q strenger eingestellt werden. Das hat aber gleichzeitig eine Verringerung der *Effizienz* zur Folge (Abschnitt 5.4.1). Die gewählten Kriterien für die Überführung des pixelbasierten Klassifikationsergebnisses auf ein GIS-Objekt einfach nur strenger einzustellen, ist daher nicht sinnvoll, was auch Abbildung 43 belegt. Abbildung 43 zeigt ein Beispiel dafür, dass es bei der pixelbasierten Klassifikation

oft zur Verwechslung von unbewachsenem Ackerland und Industrie kommt (durch einen Pfeil gekennzeichnet). Dieses Problem wurde bereits in Kapitel 4.8.1 erörtert. Ein einfaches Anpassen der Parameter bei der Überführung des pixelbasierten Klassifikationsergebnisses auf ein GIS-Objekt ist daher nicht sinnvoll und man erreicht lediglich eine sehr gute Anpassung des Problems auf eine einzelne Szene. Daher ist die Wahl anderer bzw. zusätzlicher Kriterien für die Überführung der pixelbasierten Klassifikation auf ein GIS-Objekt nötig, die im Ausblick in Kapitel 6 diskutiert werden. In Anbetracht der geringen Fehlerquote, sind die hier verwendeten Kriterien jedoch als zunächst ausreichend zu betrachten.

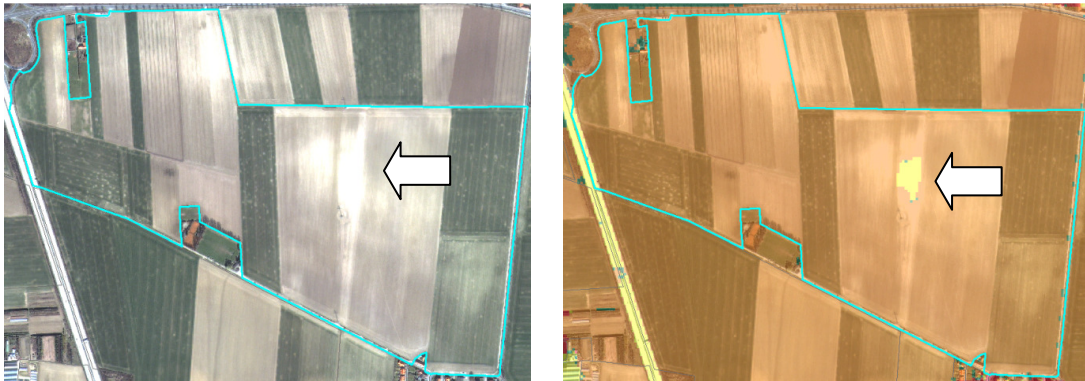


Abbildung 43: Korrektes GIS-Ackerlandobjekt, bei dem teilweise fälschlicherweise Industrie klassifiziert wurde (Testszene: Hildesheim). Links: RGB-Bild, rechts: Ergebnis der pixelbasierten Klassifizierung (braun: Acker-/Grünland, gelb: Industrie, rot: Siedlung, grün: Wald).

Bei der detaillierteren Analyse der *False Positive* in der Testszene Halberstadt unter Verwendung von ATKIS, in der das schlechteste Ergebnis bezüglich der *TG a posteriori* erreicht werden konnte, wurden 13 von insgesamt 131 fehlerhaften GIS-Objekten nicht detektiert. Bei den 13 nicht gefundenen fehlerhaften GIS-Objekten handelte es sich um GIS-Ackerland- oder Grünlandobjekte, bei denen Teilflächen fremder Nutzung (versiegelte Flächen) vorlagen. Diese bebauten Flächen konnten jedoch ähnlich wie in Abbildung 42 auf Grund der geringen Größe weder als kompakte Fehler noch über den Schwellwert s_q für fehlerhafte Pixel ermittelt werden. Es gilt dieselbe Diskussion, die bereits oben geführt wurde.

Die *Effizienz* des Systems ist bis auf die der Szene Halberstadt sehr hoch und liegt bei über 80% unter Betrachtung der Anzahl der Objekte und bei über 95% bei der Betrachtung der Fläche. In der Szene Halberstadt haben besonders viele kleine Objekte zu einer hohen Rate an verworfenen GIS-Objekten geführt, was die *Effizienz* des Systems vermindert.

Abschließend ist zusammenzufassen zu sagen, dass mit dem ersten Analyseschritt, wie in Abschnitt 4.3 beschrieben, sehr gute Ergebnisse erreicht werden und daher für die Gesamtmethode zur Verifikation von GIS-Ackerland- und Grünlandobjekten einsetzbar ist. Potentiale zur Verbesserung sind bei der Überführung des pixelbasierten Klassifikationsverfahrens auf die GIS-Objekte vorhanden. Es empfiehlt sich ein Schwellwert von $s_q=30\%$ oder $s_q=40\%$, wobei $s_q=30\%$ die strengere Bewertung ist.

5.5. Untersuchung des Einflusses von verschiedenen Merkmalsgruppen auf die objektbasierte Klassifikation

Kern des zweiten Analyseschrittes, der Trennung von GIS-Ackerland- und Grünlandobjekten mittels einer objektbasierten Klassifikation, sind die Merkmale der vier Merkmalsgruppen (spektral, textuell, strukturell und geometrisch), da diese die Zuordnung der Segmente zu den Klassen Acker- oder Grünland ermöglichen. In einer ersten Untersuchung soll der Einfluss der verschiedenen Merkmalsgruppen auf den Erfolg der Verifikation und, da es sich bei diesem Analyseschritt um ein neues eingeführtes Verfahren handelt, diesmal auch auf den Erfolg der Klassifikation explizit getestet werden. Daher wird unter Verwendung verschiedener Kombinationen von Merkmalsgruppen ein Testdatensatz evaluiert. Hierzu wurde das ausgedünnte Schlagkataster der Testszene Halberstadt verwendet, da:

- es sich um ein Schlagkataster handelt und eine Segmentierung in Schläge nicht notwendig ist und
- eine Vielzahl von Ackerland- und Grünlandobjekten vorliegt.

Für die Analyse wurden insgesamt 8 Tests mit jeweils unterschiedlichen Kombinationen von Merkmalsgruppen durchgeführt. Ein Überblick ist in Tabelle 28 gegeben.

Test	Abkürzung	Verwendete Merkmalsgruppen			
		Spektrale	Texturelle	Strukturelle	Geometrisch
1	x_{spe}	X			
2	x_{tex}		X		
3	x_{stru}			X	
4	x_{tex_stru}		X	X	
5	x_{spe_stru}	X		X	
6	x_{spe_tex}	X	X		
7	$x_{spe_tex_stru}$	X	X	X	
8	$x_{spe_tex_stru_form}$	X	X	X	X

Tabelle 28: Überblick über die verschiedenen Tests zur Analyse des Einflusses unterschiedlicher Merkmalsgruppen auf die objektbasierte Klassifikation des zweiten Analyseschrittes.

5.5.1. Evaluierung der Klassifikationsergebnisse

Zunächst werden die Ergebnisse der Evaluierung bezüglich der Klassifikation ausgewertet. Abbildung 44 zeigt die Darstellung der Klassifikationsergebnisse für alle in Tabelle 28 aufgeführten Szenarien. Zunächst lässt sich allgemein feststellen, dass die Ergebnisse für die Klassifikation der Klasse Ackerland immer besser als die Ergebnisse der Klassifikation der Klasse Grünland sind. Die *overall accuracy* erreicht immer 85%, auch wenn die *producer's accuracy* von Grünland teilweise sehr niedrig ist (weniger als 65%), da mehr sicher klassifizierbare Ackerlandobjekte als nicht sicher klassifizierbare Grünlandobjekte im Datensatz vorhanden sind.

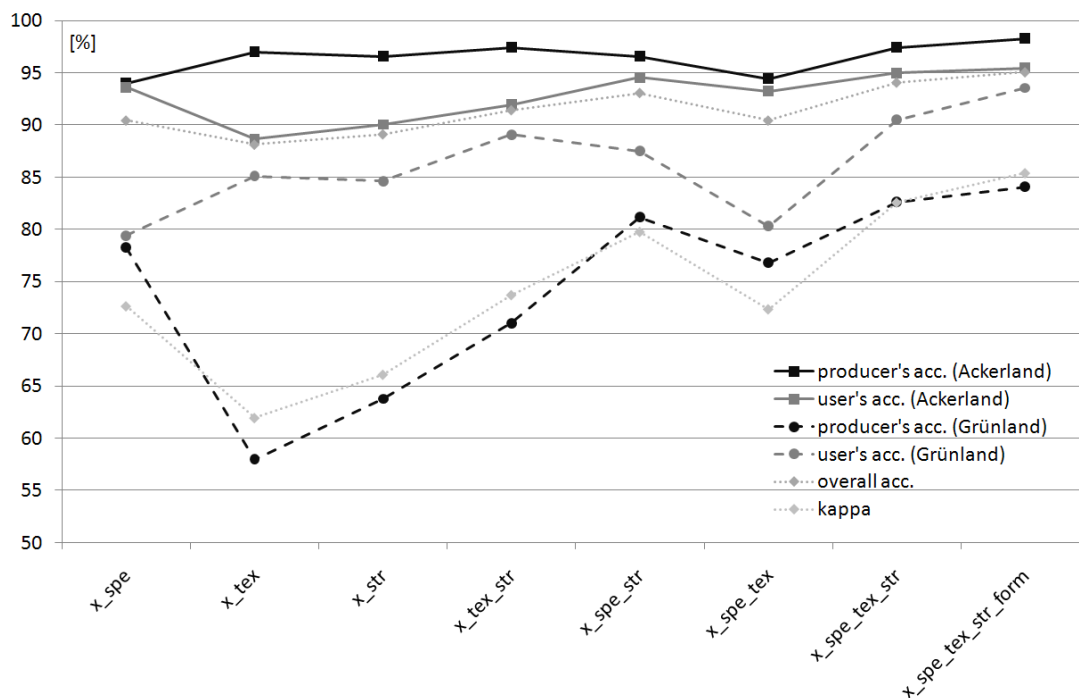


Abbildung 44: Evaluierungsergebnisse für die Klassifikation bzgl. der Untersuchung des Einflusses der unterschiedlichen Merkmalsgruppen auf die objektbasierte Klassifikation des zweiten Analyseschrittes.

Zunächst werden nur die Ergebnisse der ersten drei Tests betrachtet, bei denen jeweils nur eine Merkmalsgruppe verwendet wird. Das beste Ergebnis konnte bei Betrachtung des Kappa-Koeffizienten mit den spektralen Merkmalen erreicht werden, das schlechteste mit den texturellen Merkmalen. Der Kappa-Koeffizient liegt immer bei über 60%, was nach Fleiss (1983) als gut eingeordnet wird. Die Aussage, dass die texturellen Merkmale den geringsten Einfluss auf den Erfolg der Klassifikation haben, kann auch auf die *producer's accuracy* von Grünland und der *user's accuracy* von Ackerland bezogen werden. Die Erkenntnis findet sich auch in den nächsten drei Tests bestätigt, die die Ergebnisse von jeweils zwei verwendeten Merkmalsgruppen zeigen. Das beste Ergebnis konnte erreicht werden, wenn die texturellen Merkmale weggelassen werden. Die Ergebnisse für die Kombination zwischen texturellen Merkmalen mit jeweils

spektralen und strukturellen Merkmalen sind vergleichbar; der Kappa-Koeffizient liegt jeweils bei ca. 73%. Werden die Merkmale der spektralen, texturellen und strukturellen Merkmalsgruppen kombiniert, steigen auch *overall accuracy*, *producer's accuracy* von Grünland und Kappa-Koeffizient auf über 80%. Die Werte für die *producer's accuracy* von Ackerland bei der Kombination von Merkmalen aller drei Merkmalsgruppen sind jedoch nur geringfügig höher als bei der Kombination von nur spektralen und strukturellen Merkmalen, was ein weiterer Hinweis darauf ist, dass zumindest in der Testszene Halberstadt die texturellen Merkmale den geringsten Einfluss auf den Erfolg der Klassifikation haben. Bei der zusätzlichen Hinzunahme von geometrischen Merkmalen zu der Klassifikation kann die Genauigkeit der Ergebnisse weiter gesteigert werden. Während die Genauigkeitssteigerung bezüglich der Klasse Ackerland eher gering ausfällt, ist sie bei Grünland stärker. Das zusätzliche Verwenden von geometrischen Merkmalen scheint also besonders bei der Verifikation von Grünland von Vorteil zu sein, da viele sonst schwer zu klassifizierende Grünlandobjekte sehr markante Formen besitzen (langgestreckte Objekte). Die 75% Marke bei Betrachtung des Kappa-Koeffizienten (nach Fleiss (1983) ausgezeichnet) wird immer dann erreicht, wenn spektrale und strukturelle Merkmale verwendet werden.

5.5.2. Evaluierung der Verifikationsergebnisse

Nachdem die Evaluierungsergebnisse der Klassifikation näher betrachtet wurden, erfolgt nun die Analyse der Evaluierungsergebnisse der Verifikation. Die Ergebnisse bezüglich der Größen *TG a priori* und *TG a posteriori* sind graphisch in Abbildung 45 dargestellt. Während die *TG a priori* von Grünland bei 71,9% und damit deutlich unter der geforderten Genauigkeit von 95% liegt, beträgt die *TG a priori* von Ackerland 97,7%. Die *TG a priori* der Vereinigungsmenge aller Objekte der Klassen Acker- und Grünland liegt bei 90,1% und damit ebenfalls unter der geforderten Marke von 95%. Sowohl für die GIS-Ackerland- also auch für die GIS-Grünlandobjekte und der Vereinigungsmenge aller GIS- Ackerland- und Grünlandobjekte liegt die *TG a posteriori* über der geforderten Marke von 95%. Die *TG a posteriori* ist bei den Ergebnissen, die mit nur einer Merkmalsgruppe berechnet wurden, am geringsten. Die größten Schwankungen der *TG a posteriori* kann bei der Klasse Grünland beobachtet werden. Dort fällt besonders auf, dass die *TG a posteriori* beim Test mit nur spektralen und texturellen (also ohne strukturelle) Merkmalen genauso hoch ist wie bei der Verwendung von Merkmalen aus allen Merkmalsgruppen (ausgenommen sind die Tests, die mit nur einer Merkmalsgruppe berechnet wurden). Zusammenfassend kann aber gesagt werden, dass die geforderte *TG a posteriori* von 95% bei dieser Testszene bei allen getesteten Konfigurationen erreicht werden konnte.

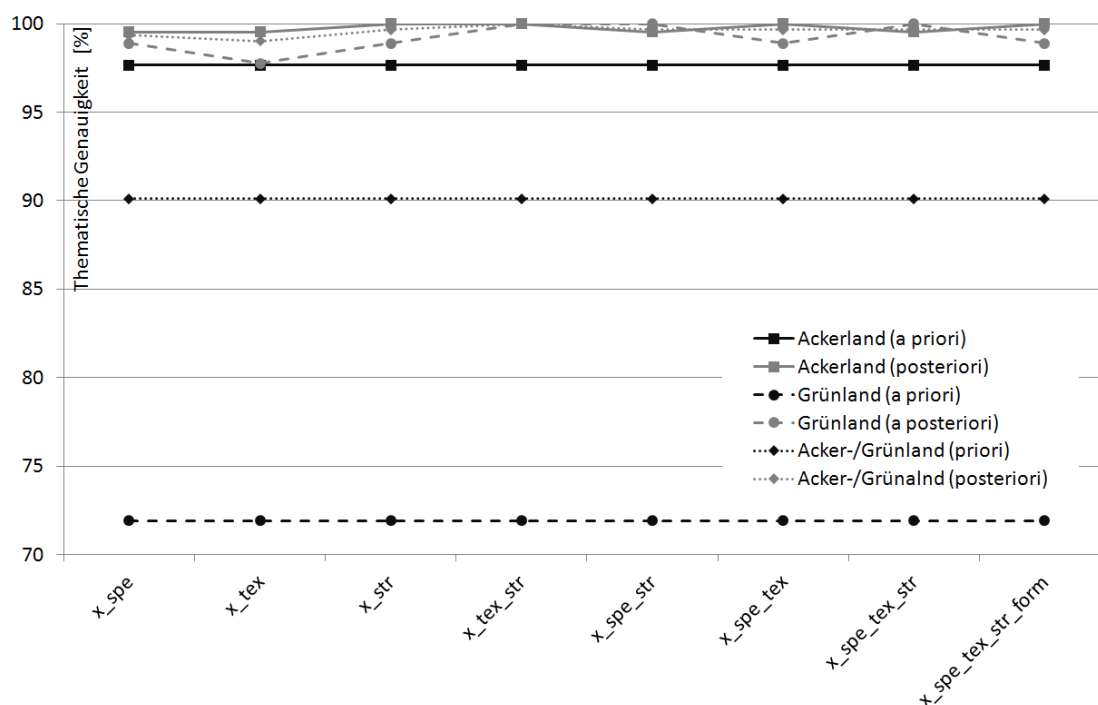


Abbildung 45: Evaluierungsergebnisse der Verifikation bzgl. der Untersuchung des Einflusses der unterschiedlichen Merkmalsgruppen auf die objektbasierte Klassifikation, hier: Betrachtung der *TG a priori* und *TG a posteriori*.

Die *Effizienz*, die Anzahl der *False Positives* und die *overall accuracy* sind in Abbildung 46 dargestellt. Während die *Effizienz* der Klasse Ackerland immer sehr hoch ist (über 90%), ist die *Effizienz* der Klasse Grünland eher gering und liegt zwischen 42,7% und 60,7%. Da wesentlich mehr Ackerland- als Grünlandobjekte vorhanden sind, ist die *Effizienz* der Vereinigungsmenge aller Ackerland- und Grünlandobjekte hoch und liegt zwischen 79,9% und 85,8%. Die Vereinigungsmenge aller Ackerland- und Grünlandobjekte wird betrachtet, weil nur dadurch der Arbeitsumfang des menschlichen Bearbeiters für die manuelle Nachbearbeitung wirklich abgeschätzt werden kann. Die *Effizienz* ist am höchsten, wenn Merkmale aus mindestens drei Merkmalsgruppen verwendet werden.

Der Verlauf der *overall accuracy* in Abbildung 46 ist vergleichbar mit dem der *Effizienz*, wobei sie auf einem höheren Niveau liegt. Auf eine Darstellung der *error detection rates* wird zu Gunsten einer besseren Übersicht in der Abbildung 46 verzichtet. Sie beträgt bei Ackerland immer über 80%, bei Grünland über 92% und bei der Vereinigungsmenge bei über 90%, was sehr zufriedenstellende Ergebnisse sind. Diese Rate ist ausreichend, damit die *TG a posteriori* immer bei über 95% liegt.

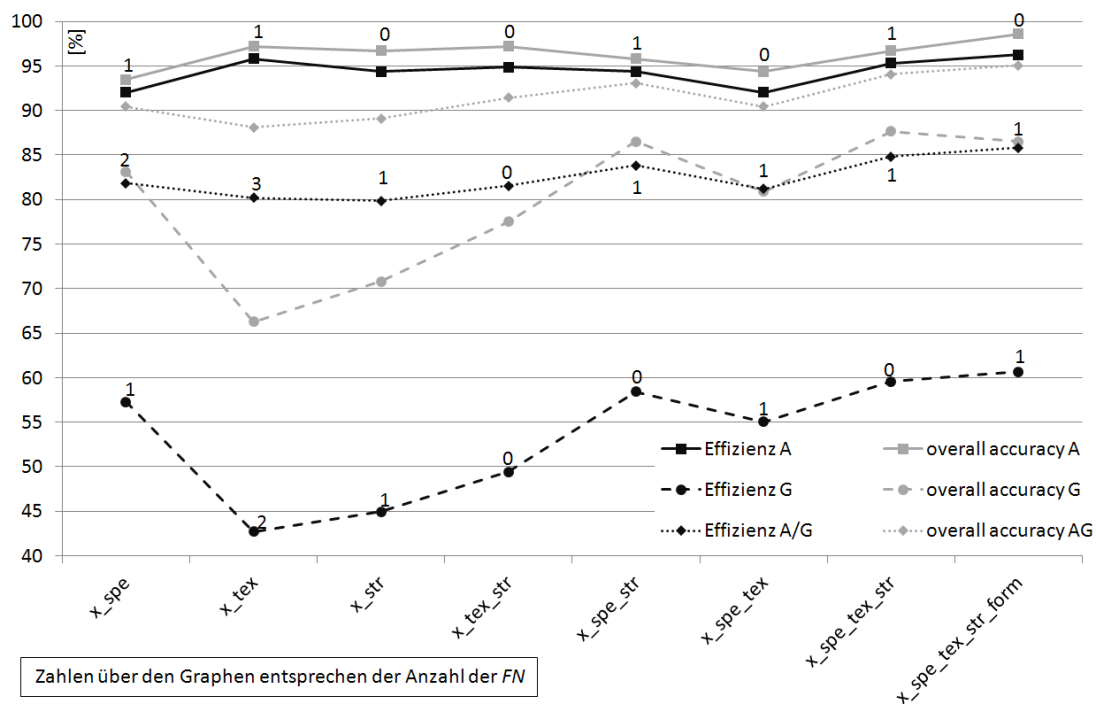


Abbildung 46: Evaluierungsergebnisse der Verifikation bzgl. der Untersuchung des Einflusses der unterschiedlichen Merkmalsgruppen auf die objektbasierte Klassifikation, hier: Betrachtung der *Effizienz*, *overall accuracy* und Anzahl der *False Positives* (Zahlen über den Graphen) für A... Ackerland, G... Grünland und A/G... der Vereinigungsmenge aus Acker- und Grünland.

Betrachtet man nur die Anzahl der *False Positives*, so konnte das beste Ergebnis mit den Merkmalen aus den Merkmalsgruppen der textuellen und strukturellen Merkmale erreicht werden. Die Anzahl der *False Positives* beträgt hier Null. Ebenfalls gute Ergebnisse mit nur einem verbleibenden Fehler im GIS konnten für die Kombinationen

- Nur strukturelle Merkmale (ein *False Positives*, Objektnr. 11888),
- Spektrale und strukturelle Merkmale (ein *False Positives*, Objektnr. 12477),
- Spektrale und textuelle Merkmale (ein *False Positives*, Objektnr. 11867),
- Spektrale, textuelle und strukturelle Merkmale (ein *False Positives*, Objektnr. 12477) und
- Merkmale aus allen Merkmalsgruppen (ein *False Positives*, Objektnr. 11867)

erreicht werden. Die *False-Positive*-Objekte sind in Abbildung 47 zu sehen. Das linke Objekt wird in der GIS-Datenbank als Grünlandobjekt geführt, tatsächlich handelt es sich jedoch um ein Ackerlandobjekt, das auch mit einer anderen Teilnutzung bedeckt ist. Beim mittleren Objekt handelt es sich um ein bewachsenes GIS-Ackerlandobjekt, das im GIS als Grünlandobjekt geführt wird. Dieses Objekt ist jedoch nicht groß genug, damit Strukturen darin detektiert werden können. Der Algorithmus bewertet es daher als Grünlandobjekt, da die spektralen Merkmale dieses GIS-Objektes einem Grünlandobjekt gleichen. Beim rechten Objekt liegt der umgekehrte Fall vor, also ein Grünlandobjekt, das im GIS als Ackerlandobjekt geführt wird. Die zu

erkennenden leichten Linienstrukturen und eine für ein GIS-Grünlandobjekt nicht stark ausgeprägte Form führen zur Klassifikation des GIS-Objektes als Ackerlandobjekt.



Abbildung 47: *False Positives* mit den Objektnummern 11867 (links), 11888 (Mitte) und 12477 (rechts).

Das Evaluierungsergebnis unter der Verwendung aller Merkmalsgruppen ist in Abbildung 48 zu sehen. Es wird die selbe Farbcodierung wie in Abbildung 41 verwendet. Das einzige *False Positive* GIS-Objekt (Nr. 11867) ist kaum sichtbar und daher durch einen Pfeil hervorgehoben. Unter gleichzeitiger Betrachtung der Anzahl der *False Positives* und der *Effizienz* empfiehlt es sich, eine Kombination zu verwenden, die auf jeden Fall spektrale und strukturelle Merkmale umfasst.



Abbildung 48: Graphische Darstellung der Konfusionsmatrix der Verifikationsergebnisse unter Verwendung von Merkmalen aus allen Merkmalsgruppen für die objektbasierte Klassifikation (TP: dunkelgrün, FN: hellgrün, FP: hellrot (durch einen Pfeil gekennzeichnet), TN: dunkelrot).

5.5.3. Bewertung

Schwerpunkt der Untersuchung war zu prüfen, ob mittels der in Abschnitt 4.5 vorgestellten Merkmale unter Verwendung einer SVM eine erfolgreiche Klassifikation und Verifikation von GIS-Ackerland- und Grünlandobjekten eines Schlagkatasters möglich ist. Da hier besonders die Merkmale untersucht werden sollten, wurde ein Schlagkataster verwendet, so dass auf eine Segmentierung zunächst verzichtet werden kann. Mit den in diesem Abschnitt verwendeten Daten konnte eine erfolgreiche Klassifikation und Verifikation der GIS-Objekte durchgeführt werden. Die besten Ergebnisse werden erreicht, wenn spektrale und strukturelle Merkmale beteiligt sind. Das beste Ergebnis liegt vor, wenn neben den spektralen und strukturellen Merkmalen auch alle anderen Merkmale in die Klassifikation eingehen. Die Übertragbarkeit auf andere Datensätze bleibt zunächst offen und wird in noch folgenden Experimenten implizit untersucht.

5.6. Untersuchung des Einflusses von verschiedenen Zeitpunkten auf die objektbasierte Klassifikation

Nachdem der Einfluss der unterschiedlichen Merkmalsgruppen auf das Ergebnis der Klassifikation und Verifikation untersucht wurde, wird in diesem Abschnitt geprüft, inwieweit der Zeitpunkt der Aufnahme der Bilddaten einen Einfluss auf das Klassifikations- und Verifikationsergebnis hat. Wie bereits in Kapitel 4 diskutiert, kann sich das Erscheinungsbild von Acker- und Grünland je nach dem Stadium der Vegetation (phänologische Änderung) stark ändern. Mit diesem Test soll geprüft werden, ob in verschiedenen Stadien der Vegetation die Klassifikation/Verifikation erfolgreich durchgeführt werden kann. Dazu wird der Testdatensatz Rukieten verwendet, für den als einziger Datensatz multitemporale, hochaufgelöste und multispektrale Bilder zur Verfügung stehen.

Da das Datum der Aufnahme für diese Untersuchung nicht entscheidend ist, werden bei der Auswertung die Bezeichnungen in Tabelle 29 verwendet. Die Zuweisung von Aufnahmezeitpunkten zu den Zeitpunkten A bis F erfolgte hier zur besseren Lesbarkeit der folgenden Diagramme zur Evaluierung. Da nicht der beste Zeitpunkt für die Klassifikation/Verifikation ermittelt werden soll, sondern nur, ob eine Abhängigkeit des Ergebnisses vom Zeitpunkt der Aufnahme besteht, hätte die Zuweisung der Zeitpunkte A bis F zu den Aufnahmezeitpunkte in Tabelle 29 auch anders gewählt werden können.

Bezeichnung	Aufnahmezeitpunkt
Zeitpunkt A	25.06.2006
Zeitpunkt B	04.05.2006
Zeitpunkt C	09.06.2006
Zeitpunkt D	23.11.2005
Zeitpunkt E	17.04.2006
Zeitpunkt F	07.10.2005

Tabelle 29: Übersicht über Bezeichnung/Aufnahmezeitpunkte.

Stichprobenartige Tests haben gezeigt, dass die Einführung von geometrischen Merkmalen in dieser Testszene zu einer Verschlechterung der Ergebnisse führte. Grund dafür ist, dass in diesem Gebiet GIS-Grünlandobjekte oft sehr große sind, sowie das deren Form identisch zu den GIS-Ackerlandobjekten sind. Aus diesem Grund wurden für alle Testreihen in diesem Abschnitt nur spektrale, textuelle und strukturelle Merkmale verwendet; auf geometrische Merkmale wird verzichtet. Die Klassifikation/Verifikation wurde für jeden Zeitpunkt des Datensatzes durchgeführt, für den die Kanäle *NIR*, *R*, *G* und *B* zur Verfügung stehen.

5.6.1. Evaluierung der Klassifikationsergebnisse

In Abbildung 49 sind die Evaluierungsergebnisse der Klassifikation dargestellt. Zunächst ist anzumerken, dass die Ergebnisse, verglichen mit denen in Abschnitt 5.5, wesentlich schlechter ausfallen. Dies hat mehrere Gründe:

- Die Anzahl der Trainingsobjekte ist deutlich geringer (Ackerland: statt 234 nur 24 GIS-Objekte, Grünland: statt 70 nur 14 GIS-Objekte, s. Tabelle 15 und Tabelle 21 in Abschnitt 5.2.3), somit sind nicht alle Erscheinungsformen von Acker-/Grünland ausreichend im Trainingsdatensatz vertreten.
- Eingangs wurde bereits erklärt, dass nur Zeitpunkte verwendet wurden, an denen vier Bildkanälen zur Verfügung stehen. Trotz der Kalibrierungsprobleme des *NIR*-Kanals wurde angenommen, dass dieser Kanal und der daraus abgeleitete NDVI-Kanal wertvolle Information für die Klassifikation liefern können. Die Ergebnisse lassen diesen Schluss jedoch nicht zu.

Das schlechteste Klassifikationsergebnis wurde zum Zeitpunkt A erzielt. Erwähnenswert ist, dass die Daten der Testszene Halberstadt des Experiments im Abschnitt 5.5 zu einem ähnlichen Zeitpunkt aufgenommen wurden. Während für die Szene Halberstadt im Abschnitt 5.5 sehr gute Ergebnisse erzielt worden, kann dies mit dem Rukieten-Datensatz nicht erreicht werden. Gründe hierfür sind, dass:

- Bilddaten sowohl in ihren geometrischen also auch spektralen Merkmalen unterschiedlich sind,
- der Zeitpunkt der Aufnahme im Laufe des Tages unterschiedlich sein kann (HBS: 11:43 Uhr, RUK: unbekannt) und somit auch andere Bedingungen, z.B. für die Beleuchtung, vorgelegen haben können,
- weniger Trainingsdaten in der Testszene Rukieten zur Verfügung standen (s. oben),
- das Stadium der Entwicklung der Vegetation in unterschiedlichen Jahren voneinander abweichen kann und
- die Feldfrüchte unterschiedlich sein können und manche Feldfrüchte besser als andere klassifiziert werden können.

Weitere Beobachtungen sind, dass während die Ergebnisse für Ackerland in der Szene Rukieten zufriedenstellend sind, der Algorithmus in der Klassifikation von Grünlandobjekten keine zufriedenstellenden Ergebnisse erreichen kann. Die Beobachtungen, dass Ackerland sicherer als Grünland klassifiziert werden kann, wurden bereits im Experiment in Abschnitt 5.5 gemacht. Auf Grund der größeren Anzahl an Ackerlandobjekten ist die *overall accuracy* bei über 70%, der Kappa-Koeffizient liegt jedoch bei nur ca. 33%. Zur Erinnerung: nach Fleiss (1983) ist ein Kappa-Koeffizient zwischen 40% und 60% annehmbar, zwischen 60% und 75% als gut und darüber hinaus als ausgezeichnet anzusehen ist. Leicht bessere Ergebnisse für die Klassifikation werden bei den Zeitpunkten B und C erreicht, bei denen der Kappa-Koeffizient bei über 40% liegt. Der Kappa-Koeffizient erreicht erst für den Zeitpunkt D einen Wert von über 50%, wobei die *producer's accuracy* von Grünland immer noch nicht zufriedenstellend ist (kleiner als 50%). Nur zu den Zeitpunkten E und F werden zufriedenstellende Ergebnisse erreicht. So liegt der Kappa-Koeffizient bei 60% bzw. noch darüber (65%). Erste Vermutungen, dass die Klassifikation zu diesen sehr frühen bzw. sehr späten Zeitpunkten auf Grund eventuell nicht bewachsener Ackerlandobjekte so erfolgreich war, konnten bei Betrachtung der Bilder nicht bestätigt werden. Auch sind die strukturellen Merkmale zu diesen Zeitpunkten weniger signifikant wie z.B. verglichen mit denen zum Zeitpunkt B (dem zweitschlechtesten Ergebnis). Insgesamt wurden zwischen 5 und 10 Grünlandobjekte falsch klassifiziert, ausnahmslos handelte es sich um sehr kleine Objekte. Einige falsch klassifizierte Grünlandobjekte enthielten Spuren, zum einen auf Grund von Überlandleitungen, zum anderen auf Grund von Mäharbeiten. Es lässt sich zusammenfassend sagen, dass der Erfolg der Klassifikation vom Zeitpunkt der Aufnahme abhängt und dass nähere Untersuchungen mit einem Testdatensatz, bei dem der GIS-Datensatz mehr Ackerland- und Grünlandobjekte beinhaltet als im Schlagkataster von Rukieten und dass die Bilddaten sämtlicher Farbinformationen aus allen Kanälen *NIR*, *R*, *G*, *B* im vollem Umfang genutzt werden können. Das ist empfehlenswert.

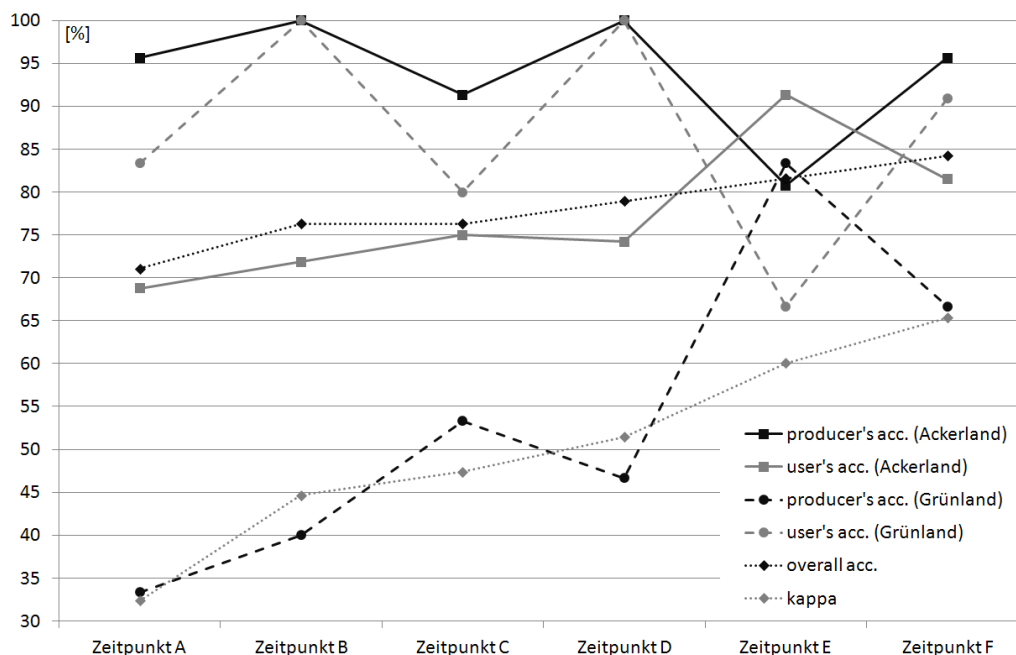


Abbildung 49: Evaluierungsergebnisse für die Klassifikation bzgl. die Untersuchung des Einflusses unterschiedlicher Zeitpunkte auf die objektbasierte Klassifikation des zweiten Analyseschrittes.

5.6.2. Evaluierung der Verifikationsergebnisse

Die Evaluierungsergebnisse der Verifikation sind in Abbildung 50 bezüglich der *TG a priori* und *TG a posteriori* und in Abbildung 51 bezüglich der *Effizienz*, der Anzahl der *False Positives* und der *overall accuracy* zusammengefasst. Wie in Abbildung 50 zu sehen ist, liegt die *TG a priori* sowohl von Ackerland (81,0%) also auch von Grünland (64,7%) und der Vereinigungsmenge der Objekte aus Acker- und Grünland (73,7%) immer deutlich unter der geforderten Marke von 95%. Mit Hilfe des Verifikationsprozesses konnten die *TG* für Ackerland auf 90,5%, für Grünland auf 88,2% bis 100% und für die Vereinigungsmenge aus Acker- und Grünland auf Werte zwischen 89,5% und 94,7% gesteigert werden. Die nach BKG (2009) geforderten Werte konnten für die Testszene Rukieten somit nicht ganz erreicht werden.

Bei Betrachtung von Abbildung 51 wird deutlich, dass bei der Verifikation von Ackerland immer zwei von vier Fehlern unentdeckt bleiben ($FP = TN$), weswegen die *Effizienz* der *overall accuracy* entspricht. Die Anzahl der *False Positives* ist in Abbildung 51 oberhalb der Graphen dargestellt. Bei den *False Positives* handelt es sich nicht immer um die gleichen Fehler, wie in Abbildung 53 deutlich wird, sondern um insgesamt vier Objekte, die immer abwechselnd nicht gefunden werden. In Abbildung 53, in der die *False Positives* von GIS-Ackerlandobjekten schwarz und die von GIS-Grünlandobjekten grau eingefärbt sind, wird deutlich, dass die *False Positives* der Objektklasse Ackerland alle sehr klein sind. Die *Effizienz* bei der Verifikation der Ackerlandobjekte ist immer sehr hoch und fällt nur bei einem Zeitpunkt unter 90%.

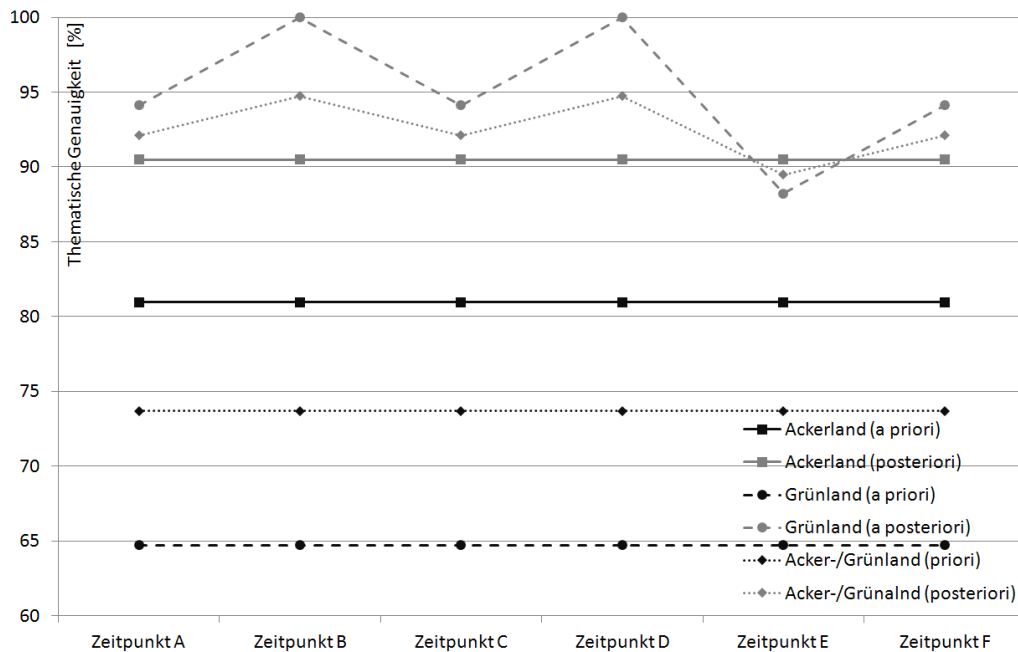


Abbildung 50: Evaluierungsergebnisse der Verifikation bzgl. der Untersuchung des Einflusses unterschiedlicher Zeitpunkte auf die objektbasierte Klassifikation, hier: Betrachtung der *TG a priori* und *TG a posteriori*.

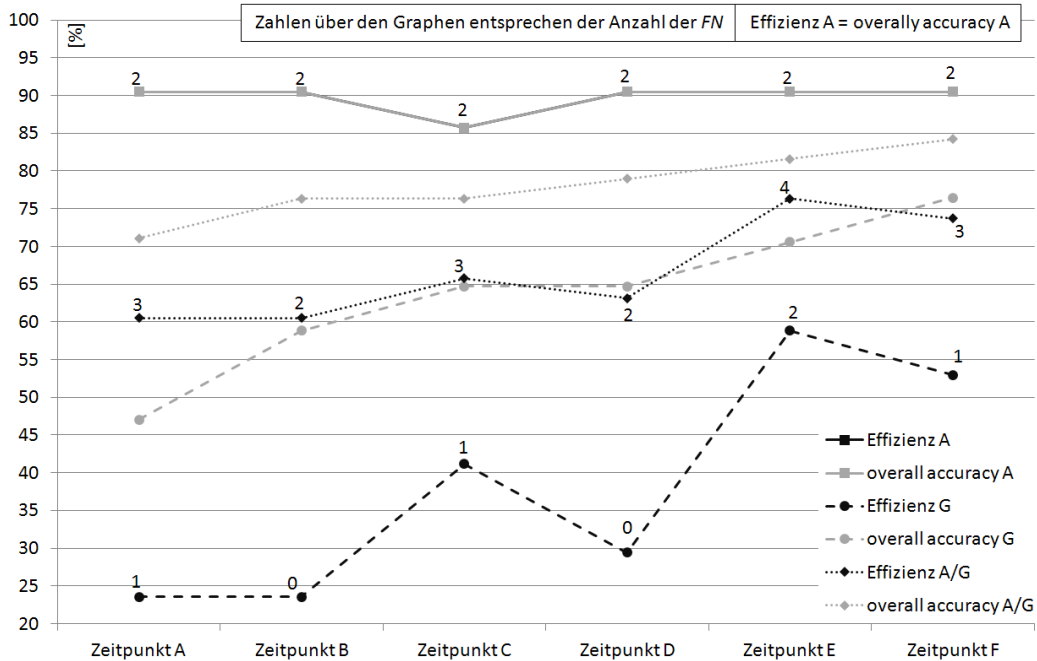


Abbildung 51: Evaluierungsergebnisse der Verifikation bzgl. der Untersuchung des Einflusses unterschiedlicher Zeitpunkte auf die objektbasierte Klassifikation, hier: Betrachtung der *Effizienz*, *overall accuracy* und Anzahl der *False Positives* (Zahlenwerte über den Graphen) für A... Ackerland, G... Grünland und A/G... der Vereinigungsmenge aus Acker- und Grünland.



Abbildung 52: *False Positive* der Klasse Grünland (um 90° gedreht).



Abbildung 53: Lage der *False Positives* (Ackerland: schwarz, Grünland: grau) bei der Evaluierung des Datensatzes RUK bzgl. der Untersuchung des Einflusses unterschiedlicher Zeitpunkte auf die objektbasierte Klassifikation.

Bei der Verifikation von Grünland gibt es zwei Zeitpunkte, bei denen alle Fehler korrekt entdeckt werden (Zeitpunkt B und D). Auch wenn die *Effizienz* generell sehr niedrig ist, so ist sie zu diesen Zeitpunkten besonders niedrig. Auf Grund des hohen Anteils von fehlerhaften GIS-Grünlandobjekten (hohe Summe aus *False Positives* und *True Negatives*) ist die *overall accuracy* hier besser geeignet, um den Erfolg des Algorithmus zu bewerten. Die *overall accuracy* kann fast immer die 60% Marke überschreiten, was ein zufriedenstellendes Ergebnis ist. Insgesamt muss der menschliche Bearbeiter aber immer abwägen, ob eine hohe *Effizienz* erreicht werden soll oder ob viele Fehler gefunden werden sollen. Dieses Problem ist auch deutlich daran zu erkennen, dass die *producer's accuracy* von Grünland in Abbildung 49 einen ähnlichen Verlauf zur *Effizienz* der dieser Klasse hat. Bei den nicht detektierten Fehlern (*False Positive*) der Objektklasse Grünland handelt es sich um insgesamt zwei GIS-Objekte. Das eine GIS-Objekt besteht im Widerspruch zur Tatsache eines Schlagkatasters zu einigen Zeitpunkten aus mehreren Schlägen, das andere Objekt ist in Abbildung 52 dargestellt. Warum dieses Objekt nicht als Ackerland klassifiziert wurde, da sowohl strukturelle als auch spektrale Merkmale für dies Klasse sprechen, kann nicht nachvollzogen werden. In Abbildung 53 ist zu erkennen, dass es sich bei den *False Positive* der Klasse Grünland nicht um kleine, sondern eher um normal große Grünlandobjekte handelt.

Die *Effizienz* und die *overall accuracy* der Vereinigungsmenge der Ackerland- und Grünlandobjekte ist zu allen Zeitpunkten zufriedenstellend. Die Anzahl der *False Positive* schwankt zwischen zwei und vier GIS-Objekten. Vier *False Positive* von insgesamt 10 im GIS enthaltenen Fehlern sind jedoch nicht zufriedenstellend. Die *TG a posteriori* für diesen Zeitpunkt (Zeitpunkt E) fällt auf unter 90% zurück. Obwohl für diesen Zeitpunkt das Ergebnis der Klassifikation sehr zufriedenstellend ist (zweitbestes Ergebnis bei der Evaluierung), können die wichtigen Objekte, die im GIS fehlerhaften Objekte, im Verifikationsprozess nicht gefunden werden. Ein gutes Klassifikationsergebnis ist daher nicht notwendigerweise Hinweis auf ein gutes Verifikationsergebnis, weil generell wenige Fehler in einem GIS enthalten sind.

Das beste Verifikationsergebnis unter Betrachtung der Vereinigungsmenge aus Ackerland- und Grünlandobjekten ist bei den Zeitpunkten erreicht, bei denen nicht mehr als zwei von 10 Fehlern unerkant bleiben. Da immer mindestens zwei *False Positive* Ackerlandobjekte vorhanden sind, fallen diese Zeitpunkte mit den Zeitpunkten zusammen, in denen es keine *False Positive* Grünlandobjekte gibt (Zeitpunkt B und D). Obwohl die *Effizienz* zu diesen Zeitpunkten für die Grünlandobjekte sehr gering ist, ist die *Effizienz* der Vereinigungsmenge der Ackerland- und Grünlandobjekte mit über 60% zufriedenstellend. Dies bedeutet, dass weniger als die Hälfte der Objekte vom menschlichen Bearbeiter manuell kontrolliert werden müssen.

5.6.3. Bewertung

Das Ziel der Untersuchung war zu evaluieren, ob der Erfolg von Klassifikation und Verifikation zeitlich von den Zeitpunkten, zu denen die Bilddaten erfasst wurden, abhängt. Die Untersuchung konnte eine zeitliche Abhängigkeit feststellen, wobei für die Klassifikation andere Zeitpunkte besser geeignet waren als für die Verifikation. Eine weitere Schlussfolgerung ist daher, dass ein gutes Klassifikationsergebnis noch kein gutes Verifikationsergebnis bedeutet. Die in (BKG, 2009) geforderte *TG a posteriori* konnten für die Testszene Rukieten nicht ganz erreicht werden. Mögliche Gründe dafür sind, dass zum einen der Datensatz nur wenige Ackerland- und Grünlandobjekte umfasst und dadurch bedingt die Anzahl der Trainingsgebiete, und zum anderen der wichtige *NIR*-Kanal und der daraus abgeleitete *NDVI* auf Grund von Kalibrierungsproblemen nur wenige zusätzliche Informationen liefern konnten. Weitere Tests der Abhängigkeit des Erfolges der Klassifikation/Verifikation des objektbasierten Klassifikationsverfahrens zur Trennung von Acker- und Grünland vom Zeitpunkt der Aufnahme sind mit einem größeren Datensatz und multispektralen Bildern ratsam.

5.7. Evaluation des zweiten Analyseschrittes

In diesem Abschnitt soll nun der gesamte zweite Analyseschritt des Verfahrens, die objektbasierte Klassifikation zur Trennung von Acker- und Grünland inklusive Segmentierung, auf tatsächliche GIS-Daten angewendet werden. Für diese Evaluation werden dieselben Testszenen wie in Abschnitt 5.4 verwendet, also die Testszenen Halberstadt (mit Schlagkataster und ATKIS) sowie Hildesheim und Weiterstadt (jeweils mit ATKIS). Beim Schlagkataster der Testszene Halberstadt wird das Verfahren des zweiten Analyseschrittes nur ohne Benutzung der Segmentierung ausgeführt, da die Schläge in diesem GIS bereits vorliegen und ihre Grenzen identisch mit den GIS-Objektgrenzen sind. Die Benutzung der Segmentierung ist nur bei dem Vorhandensein von mehreren Schlägen in einem GIS-Objekt notwendig. Die Anwendung der Segmentierung sollte jedoch in Abhängigkeit von der Größe des GIS-Objektes betrachtet werden, denn die Anwendung der

Segmentierung kann bei kleinen GIS-Objekten mit nur einem Schlag nachteilig sein. Die Nachteile entstehen dadurch, dass durch die Segmentierung und das Auftreten von möglichen kleinen Segmenten im Segmentierungsergebnis die Fläche zur Analyse des Objektes unnötig eingeschränkt wird. Aus diesem Grund werden für jede Testszene drei Testszenarien gerechnet: ohne Segmentierung, mit Segmentierung sowie ein Kombinationsszenario. Erfahrungen zeigen, dass, wenn GIS-Objekte kleiner als die Mindestkartierfläche (z.B. 1 ha bei ATKIS) sind, eine Segmentierung nicht sinnvoll ist. Diese GIS-Objekte bestehen i.d.R. nur aus einem Schlag. Wenn die GIS-Objekte also kleiner oder gleich der Mindestkartierfläche sind, wird im Kombinationsszenario das Klassifikationsergebnis ohne Anwendung der Segmentierung bestimmt. Sind die GIS-Objekte größer als die Mindestkartierfläche, wird das Ergebnis mit Benutzung der Segmentierung berechnet. Es ergeben sich insgesamt 10 Testszenarien, die in Tabelle 30 zusammengefasst sind. Bei der Evaluierung werden nur GIS-Objekte der Klasse Acker- und Grünland betrachtet, die tatsächlich einer dieser beiden Klassen angehören, also in der Referenz als GIS-Ackerland- oder Grünlandobjekte geführt werden, da das Verfahren des zweiten Analyseschrittes nur in der Lage ist, zwischen diesen beiden Klassen zu unterscheiden. Auf eine explizite Untersuchung der Segmentierungsergebnisse wird verzichtet. Visuelle Vergleiche zeigten den Erfolg der Segmentierung (auch in Abbildung 11, Abbildung 24, Abbildung 25 und Abbildung 56 zu sehen). Implizit wird die Segmentierung mit diesem Versuch getestet.

Test-szenario	Ort			Konfiguration			
	HBS	WS	HI	Schlagkataster (ohne Segmentierung)	ATKIS ohne Segmentierung	ATKIS mit Segmentierung	ATKIS Kombi.
1	X			X			
2	X				X		
3	X					X	
4	X						X
5		X			X		
6		X				X	
7		X					X
8			X		X		
9			X			X	
10			X				X

Tabelle 30: Testszenarien für die Evaluation des zweiten Analyseschrittes.

5.7.1. Evaluierung der Klassifikationsergebnisse

Die Evaluationsergebnisse der Klassifikation sind in Abbildung 54 und Abbildung 57 zusammengefasst. Zur Erinnerung: wenn der zweite Analyseschritt mit Segmentierung ausgeführt wird, können GIS-Objekte nicht nur der Klasse Ackerland oder Grünland, sondern auch einer Zurückweisungsklasse zugeordnet werden (Abschnitt 4.6.). Da eine solche Zurückweisungsklasse im GIS nicht existiert, werden bei der Evaluierung der Klassifikationsergebnisse zwei Diagramme erstellt. Im ersten Diagramm werden nur GIS-Objekte berücksichtigt, die vom System entweder der Klasse Acker- oder Grünland zugeordnet werden (Abbildung 54). Im zweiten Diagramm werden auch die Ergebnisse berücksichtigt, bei dem GIS-Objekte der Zurückweisungs-klasse zugeordnet werden (Abbildung 57).

In Abbildung 54 ist zu sehen, dass sowohl die *producer's accuracy* als auch die *user's accuracy* der Klasse Ackerland in allen Testszenarien zufriedenstellende Ergebnisse erreichen. Für die Klasse Ackerland liegt die *producer's accuracy* fast immer bei über 90% und die *user's accuracy* bei fast immer über 85%. Bei der Klasse Grünland war die Klassifikation weniger erfolgreich gegenüber der Klassifikation von Ackerland, wie dies ebenfalls bereits in vorangegangenen Tests beobachtet werden konnte. Ursache dafür ist in dieser Untersuchung, dass die Klassifikation von kleinen Objekten sich als besonders schwierig darstellt, wie später diskutiert werden wird. Und besonders bei den Grünlandobjekten handelt es sich um kleine, schwer zu klassifizierende GIS-Objekte. Während für die *user's accuracy* in der Testszene Halberstadt gute Ergebnisse mit über 85% erreicht werden konnten, liegt die *user's accuracy* in den anderen Testszenen zwischen 50% und 70%. Der Kappa-Koeffizient für Halberstadt schwankt zwischen 65% und 75%, das ist nach Fleiss (1983) ein gutes Ergebnis. Die Ursache der niedrige *user's accuracy* von Grünland in der Testszene Weiterstadt ist, dass sehr viele unbewachsene Ackerlandobjekte der Klasse Grünland zugewiesen wurden. Die geringe *user's accuracy* von Grünland bewirkt somit auch einen kleinen Kappa-Koeffizienten. Die sehr geringe *user's accuracy* von Grünland in der Szene Hildesheim ist ebenfalls auf viele unbewachsene Ackerlandobjekte, die der Grünlandklasse zugeordnet werden, zurückzuführen. Das liegt in dieser Szene vor allem daran, dass zum Zeitpunkt der Aufnahme die Eigenschaften der Grünlandobjekte sehr stark dem der

unbewachsenen Ackerlandobjekte gleichen (sehr wenig Bewuchs, keine Bewirtschaftungsspuren). Der Kappa-Koeffizient Weiterstadts für die Klassifikation ohne Segmentierung ist ausgezeichnet, für die Anwendung der Segmentierung bzw. beim Kombinationsszenario gut. Die *producer's accuracy* von Grünland ist in der Szene Weiterstadt zufriedenstellend, fällt aber in der Szene Hildesheim unter Anwendung der Segmentierung und beim Kombinationsszenario ebenfalls eher gering aus und liegt zwischen 55% und 65%. Das ist ebenfalls auf die ähnlichen Eigenschaften der Klassen Grünland und unbewachsenes Ackerland in dieser Szene zurückzuführen, weswegen nicht nur viele Ackerlandobjekte der Klasse Grünland, sondern auch umgekehrt viele Grünlandobjekte der Klasse Ackerland zugewiesen werden. Der niedrige Kappa-Koeffizient ist eine Konsequenz aus der niedrigen *producer's accuracy* und *user's accuracy* von Grünland in der Szene Hildesheim. Da der Kappa-Koeffizient bei über 40% liegt, sind die Ergebnisse noch annehmbar. Ein gutes Ergebnis für den Kappa-Koeffizienten konnte bei der Klassifikation ohne Segmentierung in der Szene Hildesheim erreicht werden. Es ist auffällig, dass die Genauigkeiten unter Benutzung der Segmentierung nur in der Testszene Halberstadt zu einer Steigerung der Ergebnisse geführt hat, was auf die besonders großen Schläge der GIS-Ackerlandobjekte in dieser Testszene zurückzuführen ist.

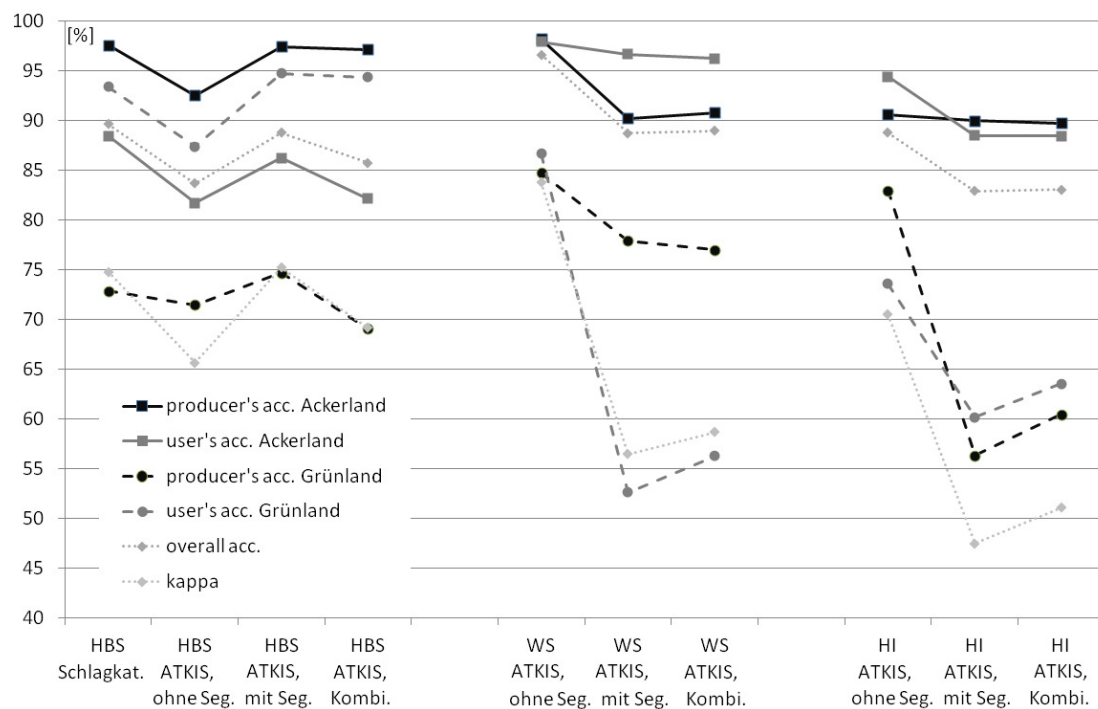


Abbildung 54: Evaluierungsergebnisse für die Klassifikation von GIS-Ackerland- und Grünlandobjekten mittels des zweiten Analyseschrittes ohne Berücksichtigung der Klassifikationsergebnisse von GIS-Objekten, die der Zurückweisungsklasse zugeordnet wurden.

Der Grund für die geringe *producer's accuracy* der Klasse Grünland in der Szene Halberstadt sind Fehlklassifikationen von GIS-Grünlandobjekten, die überstrahlt bzw. teilweise überstrahlt sind. Dies kann z.B. geschehen, wenn GIS-Objekte, wie in Abbildung 55 zu sehen, mit Sand bedeckt sind. Dann können diese GIS-Objekte leicht mit GIS-Ackerlandobjekten vertauscht werden, die kurz vor der Ernte stehen, wie sie in der Szene Halberstadt vorkommen und auch als Trainingsobjekte dienten. In der Szene Halberstadt gibt es insgesamt 16 GIS-Objekte des Schlagkatasters und ATKIS, auf die dieses Problem zu trifft. Lässt man die Objekte bei der Evaluierung der Klassifikationsergebnisse unberücksichtigt, dann erhöht sich im Fall des Schlagkatasters die *user's accuracy* von Ackerland von 88,45% auf 91,36% und die *producer's accuracy* von Grünland von 72,90% auf 80,37%. Im Fall von ATKIS (ohne Segmentierung) erhöht sich die *user's accuracy* von Ackerland von 81,74% auf 85,88% und die *producer's accuracy* von Grünland von 76,59% auf 78,97%.



Abbildung 55: Beispiele für fehlklassifizierte GIS-Objekte des Schlagkatasters der Testszene Halberstadt bei der objektbasierten Klassifikation des zweiten Analyseschrittes.

Ein Beispiel dafür, dass die Ergebnisse mit Segmentierung schlechter als ohne Segmentierung ausfallen, ist in Abbildung 56 zu sehen. Das GIS-Objekt wurde ohne Anwendung der Segmentierung korrekterweise als Ackerlandobjekt klassifiziert. Die Segmentierung hat insgesamt drei Schläge ermittelt, deren Flächengröße es zulässt, sie zu klassifizieren (rechte Darstellung in Abbildung 56). Der Schlag ganz rechts ist zu klein für die objektbasierte Klassifizierung. Zwei der drei Segmente mit einer Fläche von insgesamt 1,9202 ha wurden als Grünland klassifiziert, während das dritte Segment der Klasse Ackerland und die nicht zu klassifizierenden Segmente der Zurückweisungsklasse zugeordnet wurden. Die Flächensumme dieser Segmente erreicht 0,8144 ha. Da die Flächensumme der Segmente, die der Zurückweisungsklasse zugeordnet wurden, kleiner als die Mindestkartierfläche von 1 ha ist, wurde das GIS-Objekt fälschlicherweise als Grünlandobjekt klassifiziert. Die strukturellen Merkmale, die ansatzweise in der Abbildung zu erkennen sind, sind nicht signifikant genug, um die Segmente als Ackerland klassifizieren zu können.



Abbildung 56: Fehlerhaft klassifiziertes GIS-Objekt der Klasse Ackerland in der Szene WS (links: RGB-Bild (Dimension: 350m x 240m), rechts: Segmentierungsergebnis).

In Abbildung 57 sind die Klassifikationsergebnisse dargestellt, bei denen auch GIS-Objekte berücksichtigt werden, die der Zurückweisungsklasse zugeordnet wurden. Die Kappa-Koeffizienten sind wesentlich niedriger als bei der Betrachtung zuvor (Abbildung 54), da sehr viele Ackerland- und Grünlandobjekte der Zurückweisungsklasse zugeordnet werden (s. Tabelle 31), es diese Klasse im GIS jedoch nicht gibt. Im linken Teil der Abbildung 58 wird ein Überblick über das Klassifikationsergebnis der Szene Hildesheim unter Anwendung der Segmentierung gegeben. Hellgraue Objekte wurden der Klasse Ackerland, dunkelgraue der Klasse Grünland und rote Objekte der Zurückweisungsklasse zugeordnet. Insgesamt 110 GIS-Objekte (17,43%) werden in dieser Szene der Zurückweisungsklasse zugeordnet. Es ist gut zu sehen, dass es sich bei den meisten GIS-Objekten um eher kleine langgestreckte Objekte handelt. Diese GIS-Objekte beinhalten oft kleine Baumgruppen oder kleine Flächen anderer Nutzung, die von der Segmentierung erfolgreich bestimmt werden. Diese Segmente sind dann jedoch so klein, dass eine Klassifikation der Segmente nicht mehr möglich ist und die Mehrzahl der Segmente innerhalb eines GIS-Objektes und somit auch das GIS-Objekte der Zurückweisungsklasse zugeordnet wird. Das Gleiche gilt auch für die Szene Halberstadt, wie in Abbildung 58 auf der rechten Seite zu sehen ist. In der Szene sind jedoch gleichzeitig genug große GIS-Objekte vorhanden,

so dass dieser Effekt beim Gesamtergebnis nicht so stark wie in der Szene Hildesheim spürbar ist. Eine Kombination der Ergebnisse ohne und mit Anwendung der Segmentierung, wie es im Kombinationsszenario erfolgte, ist daher sinnvoll, was auch in Abbildung 57 an Hand der steigenden Werte der Qualitätsmaße für die Kombinationsszenarien sichtbar ist. In der Szene Weiterstadt liegen andere Bedingungen vor. Dort gibt es sehr viele GIS-Objekte mittlerer Größe, die aber unzählige kleine Schläge enthalten (Abbildung 59), die von der Segmentierung korrekt detektiert werden. Diese Schläge sind aber so klein, dass sie nicht klassifizierbar sind und somit der Zurückweisungsklasse zugeordnet werden. Gibt es mehr Segmente in einem GIS-Objekt, die so klein sind, dass sie nicht klassifizierbar sind, als Segmente, die klassifizierbar sind, wird das gesamte GIS-Objekt ebenfalls der Zurückweisungsklasse zugeschlagen. Auch hier zeigt die Kombination der Ergebnisse ohne und mit Segmentierung der Kombinationsszenarien eine Verbesserung gegenüber den Ergebnissen, die grundsätzlich mit Anwendung der Segmentierung erzeugt werden. Ansonsten sind ähnliche Tendenzen bezüglich der Klassen Acker- und Grünland, bezogen auf die *user's* und *producer's accuracy*, in Abbildung 57 wie in Abbildung 54 zu sehen. Gründe hierfür wurden bereits diskutiert. Es ist daher sogar zu empfehlen, dass der menschliche Bearbeiter vor der Verifikation kontrolliert, ob häufig Schläge innerhalb von einem GIS-Objekt zu klein für die Analyse sind. Ist dies der Fall, ist eine Analyse ohne Segmentierung zu wählen.

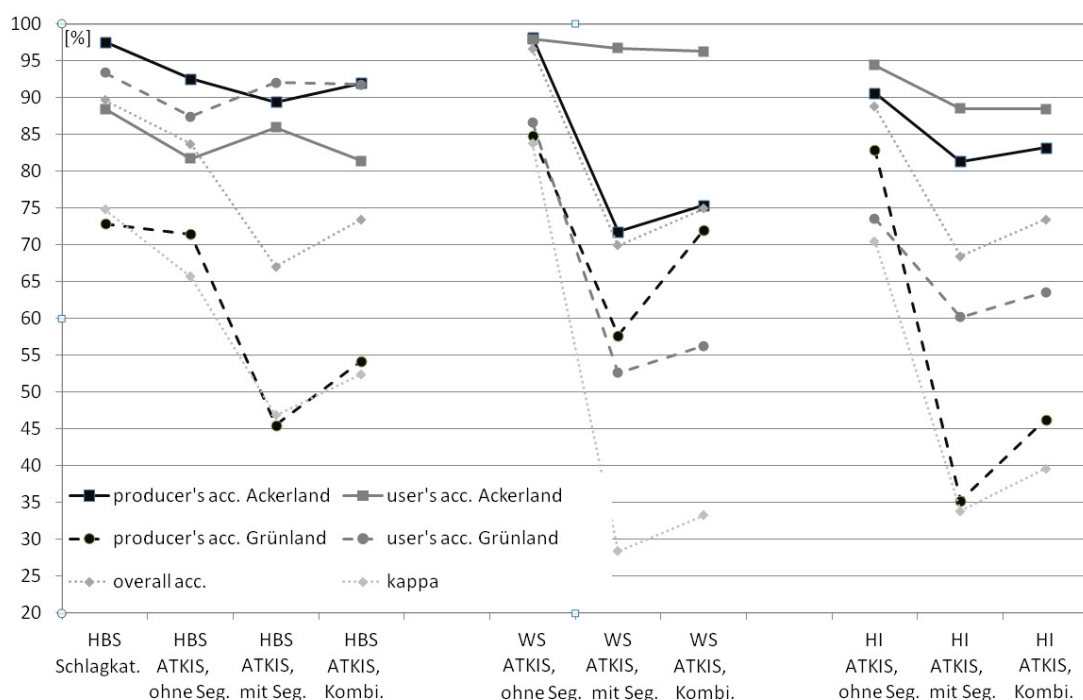


Abbildung 57: Evaluierungsergebnisse für die Klassifikation von GIS-Ackerland- und Grünlandobjekten mittels des zweiten Analyseschrittes mit Berücksichtigung der Klassifikationsergebnisse von GIS-Objekten, die der Zurückweisungsklasse zugeordnet wurden.

Testszene	Konfiguration	
	ATKIS mitSeg.	ATKIS Kombi.
Halberstadt	23,65 %	13,08 %
Weiterstadt	21,15 %	15,74 %
Hildesheim	17,43 %	11,56 %

Tabelle 31: Rate der GIS-Objekte, die der Zurückweisungsklasse zugeordnet werden (alle GIS-Objekte = 100%).



Abbildung 58: Übersicht über die Klassifikationsergebnisse der Szene HI (links) und Szene HBS (rechts) unter Verwendung der Segmentierung (hellgrau: Ackerland, dunkelgrau: Grünland, rot: Zurückweisungsklasse).

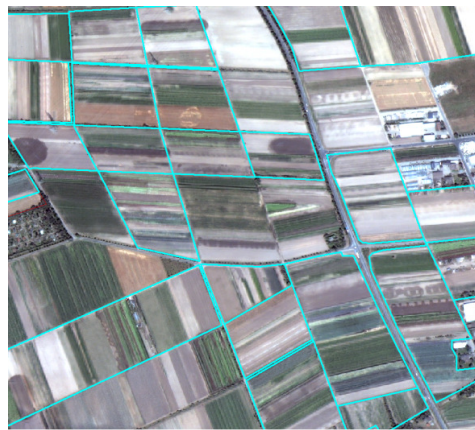


Abbildung 59: GIS-Ackerlandobjekte der Szene WS, die der Zurückweisungsklasse zugeordnet werden, da die Schläge zu klein für die objektbasierte Klassifikation sind (Objektgrenzen sind türkis hervorgehoben).

5.7.2. Evaluierung der Verifikationsergebnisse

Im Gegensatz zu den Untersuchungen bezüglich des Erfolges der Klassifikation werden bei den Untersuchungen zur Verifikation alle GIS-Objekte in die Analysen einbezogen, egal, welcher Klasse sie während der Klassifikation zugeordnet wurden, d.h. werden GIS-Objekte der Zurückweisungsklasse zugeschrieben, so werden diese vom System abgelehnt und eine Kontrolle durch den menschlichen Bearbeiter ist notwendig. Die Evaluierungsergebnisse für alle zehn Testszenarien werden in diesem Abschnitt präsentiert. Wie in Abschnitt 5.4 wird bei der Betrachtung der Ergebnisse der Verifikation nicht nur die Bewertung bezüglich der Anzahl der Objekte, sondern auch bezüglich der Fläche durchgeführt. Ziel ist das GIS zu verifizieren, so dass wie in BKG (2009) gefordert eine *TG a posteriori* von 95% bezüglich der Anzahl der Objekte und der Fläche erreicht wird.

Die Ergebnisse der *TG a priori* und *TG a posteriori*, bezogen auf die Anzahl der Objekte, sind in Abbildung 60 und, bezogen auf die Fläche, in Abbildung 61 zusammengefasst. Während in der Szene Halberstadt die gewünschte *TG a priori* für keine Klasse bezogen auf die Anzahl der Objekte, erreicht ist (Werte liegen zwischen 84% und 91%), ist für diese Szene die *TG a priori* von 95%, bezogen auf die Fläche für die Klasse Ackerland, in allen Tests erreicht. In der Szene Halberstadt fallen die *TG a priori* von Acker- und Grünland im ATKIS, bezogen auf die Anzahl der Objekte, niedriger als die *TG a priori* bezogen auf die Fläche aus. Damit wird die zuvor getroffene Beobachtung bestätigt, dass die im GIS falsch geführten Ackerland- und Grünlandobjekte kleiner als die durchschnittliche Flächengröße sind. Das Schlagkataster bildet dabei eine Ausnahme.

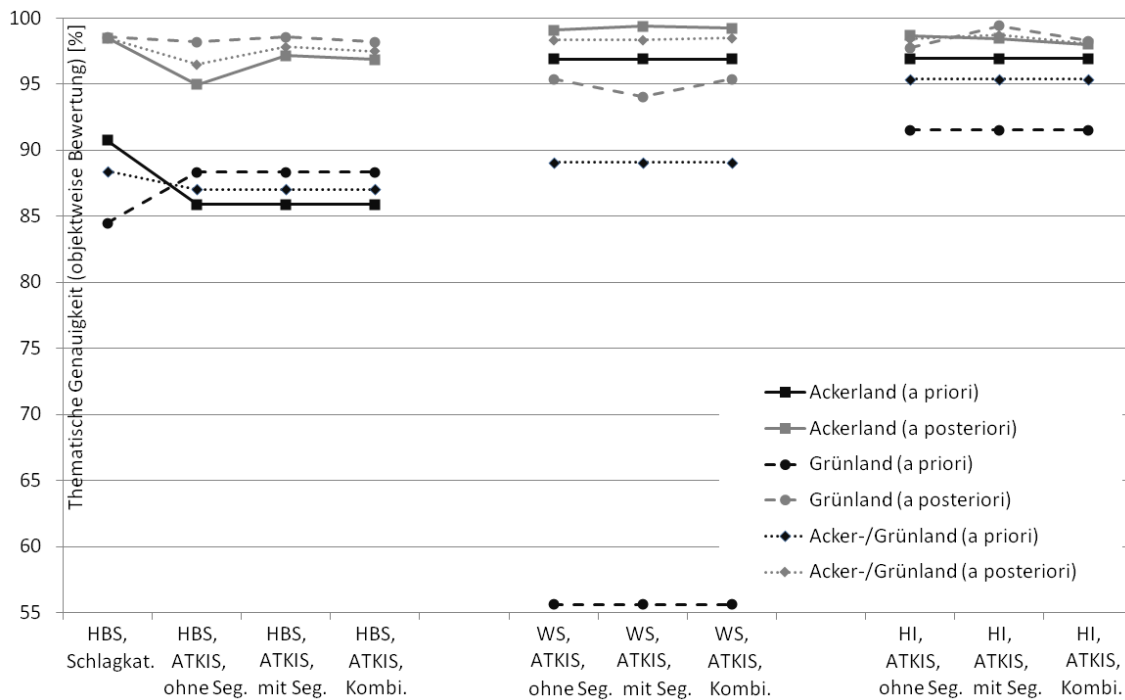


Abbildung 60: Evaluierungsergebnisse der Verifikation der Klasse Ackerland, Grünland und der Vereinigungsmenge aus Acker- und Grünlandobjekten bzgl. der *TG a priori* und *TG a posteriori* (objektweise Bewertung) des zweiten Analyseschrittes.

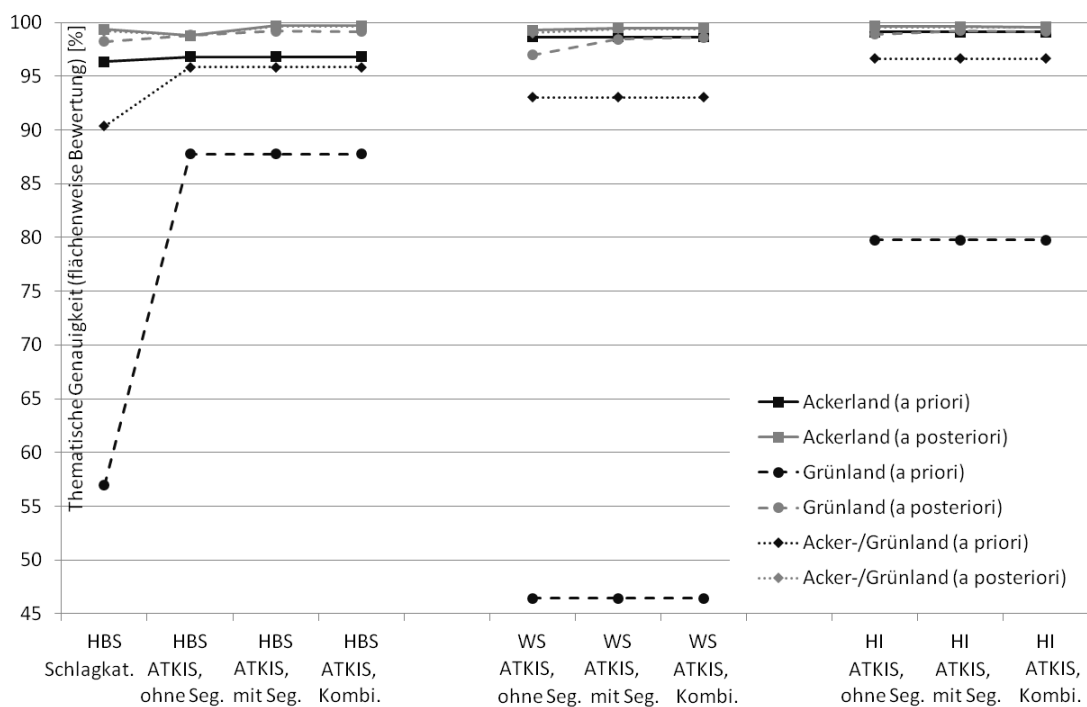


Abbildung 61: Evaluierungsergebnisse der Verifikation der Klasse Ackerland, Grünland und der Vereinigungsmenge der Acker- und Grünlandobjekte bzgl. der *TG a priori* und *TG a posteriori* (flächenweise Bewertung) des zweiten Analyseschrittes.

Nach der Verifikation konnte in allen Tests (bis auf eine Ausnahme) die gewünschte *TG a posteriori* für alle Objektarten, bezogen auf die Anzahl der Objekte und der Fläche, von 95% erreicht werden. Die Ausnahme bildet die *TG a posteriori* der Klasse Grünland, bezogen auf die Anzahl der Objekte, im Testgebiet Weiterstadt mit Anwendung der Segmentierung. Hier wurde das Ziel mit 94% knapp unterschritten. Es ist an dieser Stelle zu erwähnen, dass die *TG a priori* dieser Objektklasse bei nur 55,6% lag und damit die geringste *TG a*

priori besitzt. Das Verhalten, dass die *TG a posteriori* bei der Anwendung der Segmentierung sinkt, widerspricht dem Verhalten in den anderen Szenen, da dort die *TG a posteriori* mit Einführung der Segmentierung steigt. Eine Steigerung der *TG a posteriori* mit Einführung der Segmentierung ist zu erwarten und einfach zu erklären. Mit der Einführung der Segmentierung werden GIS-Objekte auch der Zurückweisungsklasse zugeschrieben. Objekte, die der Zurückweisungsklasse zugeschrieben werden, müssen manuell vom menschlichen Bearbeiter kontrolliert werden. Werden mehr Objekte vom System verworfen, ist es natürlich wahrscheinlich, dass auch mehr fehlerhafte GIS-Objekte vom System verworfen werden und somit die *TG a posteriori* steigt.

Bei den fehlerhaften GIS-Objekten, die in der Szene Weiterstadt ohne Anwendung der Segmentierung richtigerweise vom System verworfen wurden, aber bei der Verifikation mit Anwendung der Segmentierung unentdeckt blieben, handelt es sich um insgesamt 9 eher kleine GIS-Objekte. Ein Beispiel wurde bereits in Abbildung 56 gezeigt und diskutiert. Das im GIS als Grünland geführte Ackerlandobjekt wird ohne Anwendung der Segmentierung korrekt als Ackerlandobjekt klassifiziert. Bei Benutzung der Segmentierung wird das GIS-Objekt jedoch aus bereits beschriebenen Gründen als Grünlandobjekt klassifiziert und daher vom System akzeptiert; der Fehler im GIS bleibt unerkannt.

In Abbildung 62 sind die *error detection rates* für alle Testszenarien sowohl bezüglich der Anzahl der Objekte als auch der Fläche dargestellt. Ist die *error detection rate* bezüglich der Anzahl der Objekte höher als bezogen auf die Fläche, spricht dies für eher kleine fehlerhafte GIS-Objekte, das bestätigt zuvor gemachte Beobachtungen in Hinsicht der Größe der fehlerhaften Objekte. Wenn die *TG a priori* einer Klasse sehr hoch ist, wie die *TG a priori* der Klasse Ackerland in den Szenen Weiterstadt und Hildesheim, fällt die *error detection rate* dieser Szenarien geringer aus, da, wenn weniger Fehler im GIS vorhanden sind, i.d.R. auch weniger Fehler detektiert werden. Eine *error detection rate* von über 85% der Vereinigungsmenge von Ackerland- und Grünlandobjekten in der Szene Weiterstadt ist ein sehr zufriedenstellendes Ergebnis. Auch in der Szene Halberstadt mit einer *error detection rate* von fast immer über 80% der Vereinigungsmenge von Ackerland- und Grünlandobjekten ist ebenso sehr zufriedenstellend. Die *error detection rate* in Hildesheim fällt eher gering aus, was auf die hohe *TG a priori* und den kleineren Stichprobenumfang zurückzuführen ist.

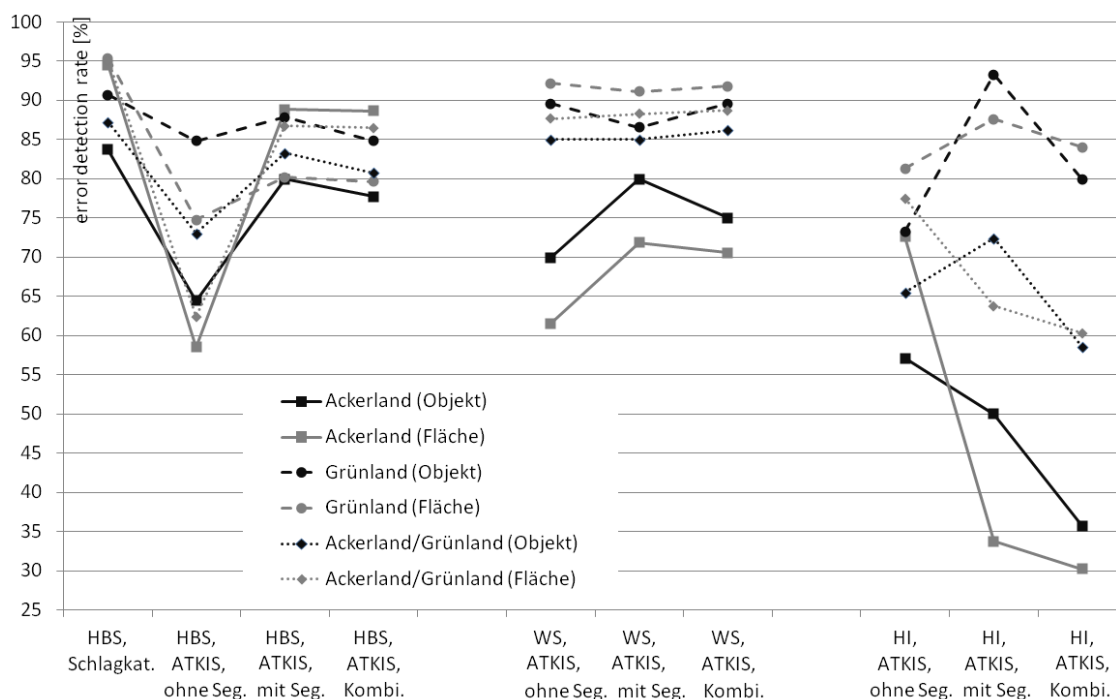


Abbildung 62: Evaluierungsergebnisse der Verifikation der Klasse Ackerland, Grünland und der Vereinigungsmenge der Acker- und Grünlandobjekte bzgl. der *error detection rate* des zweiten Analyseschrittes.

In Abbildung 63 sind die *overall accuracies* aller Testzenarien, bezogen auf die Anzahl der Objekte und der Fläche, graphisch dargestellt. Die *overall accuracies* der Klasse Grünland sind auch in diesem Versuch i.d.R. geringer als die der Klasse Ackerland. Die zuvor gemachte Aussage über die vermutliche Größe der fehlerhaften GIS-Ackerland- und Grünlandobjekte ist auch in dieser Abbildung sichtbar. Die *overall*

accuracies der Klassen Acker- und Grünland, bezogen auf die Anzahl der Objekte, sind immer kleiner als bezogen auf die Flächengröße. Während die *error detection rate* bei den Kombinationsszenarien abfällt, steigt die *overall accuracy* an. Die *overall accuracy* der Vereinigungsmenge aus Acker- und Grünland, bezogen auf die Anzahl der Objekte, liegt immer über 70%, was ein sehr zufriedenstellendes Ergebnis darstellt, da somit nur 30% der Objekte falsch zugeordnet wurden.

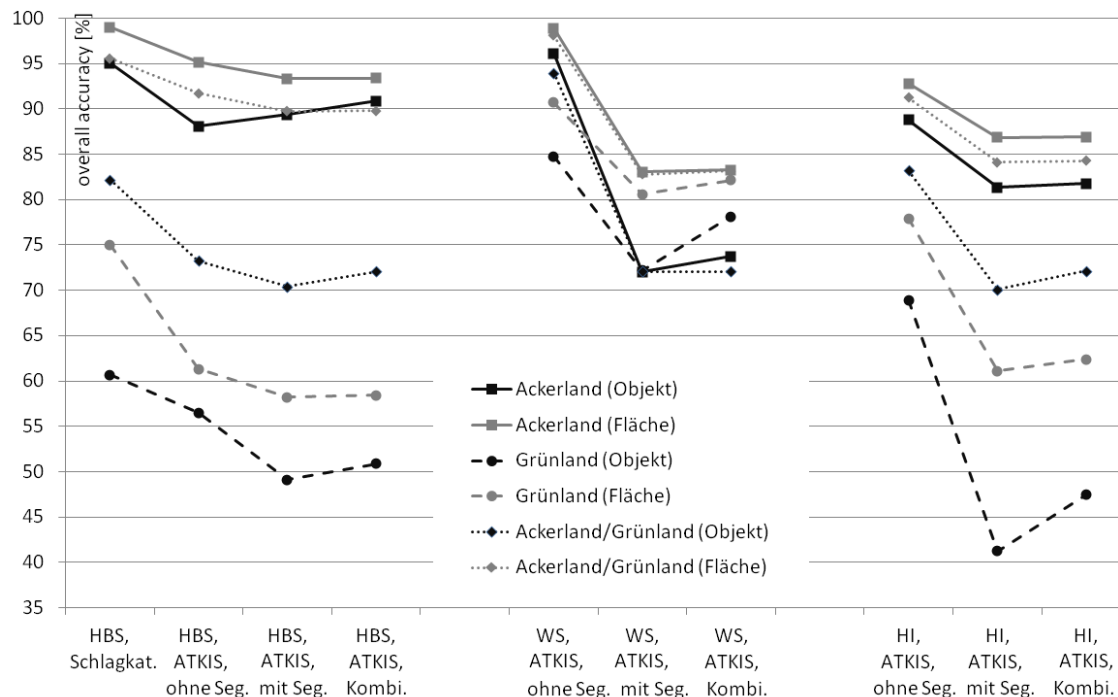


Abbildung 63: Evaluierungsergebnisse der Verifikation der Klasse Ackerland, Grünland und der Vereinigungsmenge der Acker- und Grünlandobjekte bzgl. der *overall accuracy* des zweiten Analyseschrittes.

In Abbildung 64 und Abbildung 65⁷ sind die Verifikationsergebnisse unter dem Aspekt der *Effizienz* dargestellt. Die *Effizienz* in allen drei Testszenen lässt bei der Verifikation mit Segmentierung gegenüber einer Verifikation ohne Segmentierung nach. Während die *Effizienz* in den Szenen Halberstadt und Weiterstadt nur leicht abfällt, ist der Unterschied in der Szene Hildesheim wesentlich. Ursache ist, dass durch die Einführung der Zurückweisungsklasse bei der Verifikation mit Segmentierung mehr Objekte vom System verworfen werden, weil sie der Zurückweisungsklasse zugeordnet wurden, wie oben diskutiert. Wie bereits erwähnt, müssen alle GIS-Objekte, die der Zurückweisungsklasse zugeordnet werden, vom menschlichen Bearbeiter manuell überprüft werden, damit sinkt die *Effizienz*. Ähnlich wie bei der *overall accuracy* der Kombinationsszenarien steigt auch die *Effizienz* an, besonders bezogen auf die Anzahl der Objekte. Auch das spricht für die bereits genannte Beobachtung, dass besonders viele kleine Objekte vom System bei der Verwendung der Segmentierung unberechtigter Weise abgelehnt werden. Generell kann auch in dieser Untersuchung festgestellt werden, dass die *Effizienz* bei der Verifikation von GIS-Ackerlandobjekten immer größer als die von GIS-Grünlandobjekten ist. Während die *Effizienz* für Ackerland immer größer als 80% ist (Ausnahme bildet die Testszene Weiterstadt ohne Anwendung der Segmentierung, bezogen auf die Anzahl der Objekte mit einer *Effizienz* von nur 70%), liegt die *Effizienz* von Grünland zwischen 15% und 55%, bezogen auf die Anzahl der Objekte und zwischen 35% und 75%, bezogen auf die Fläche. Betrachtet man die Vereinigungsmenge der Ackerland- und Grünlandobjekte, was dem tatsächlichen Aufwand des menschlichen Bearbeiters beim Verifikationsprozess entspricht, dann liegt die *Effizienz*, bezogen auf die Anzahl der Objekte, fast immer über 55% (Ausnahme ist die Szene Halberstadt, wo die *Effizienz* auf bis zu 48,5% abfällt). Für den menschlichen Bearbeiter entsteht damit fast immer eine Zeitersparnis um die Hälfte, verglichen zum Arbeitsaufwand ohne Verwendung des Verifikationssystems.

⁷ In Abbildung 64 ist die *Effizienz* von ‚Ackerland‘ und ‚Ackerland (> 0,5ha)‘ für Hildesheim identisch, so dass hier nur fünf statt sechs Graphen zu sehen sind. In Abbildung 65 sind die *Effizienz* von ‚Ackerland‘ und ‚Ackerland mit einer Fläche > 0,5ha‘ sowie ‚Vereinigungsmenge aus Acker-/Grünland‘ und ‚Vereinigungsmenge aus Acker-/Grünland mit einer Fläche > 0,5ha‘ fast immer identisch, weswegen die Anzahl der Graphen in der Abbildung kleiner ist als die, die in der Legende zu sehen sind.

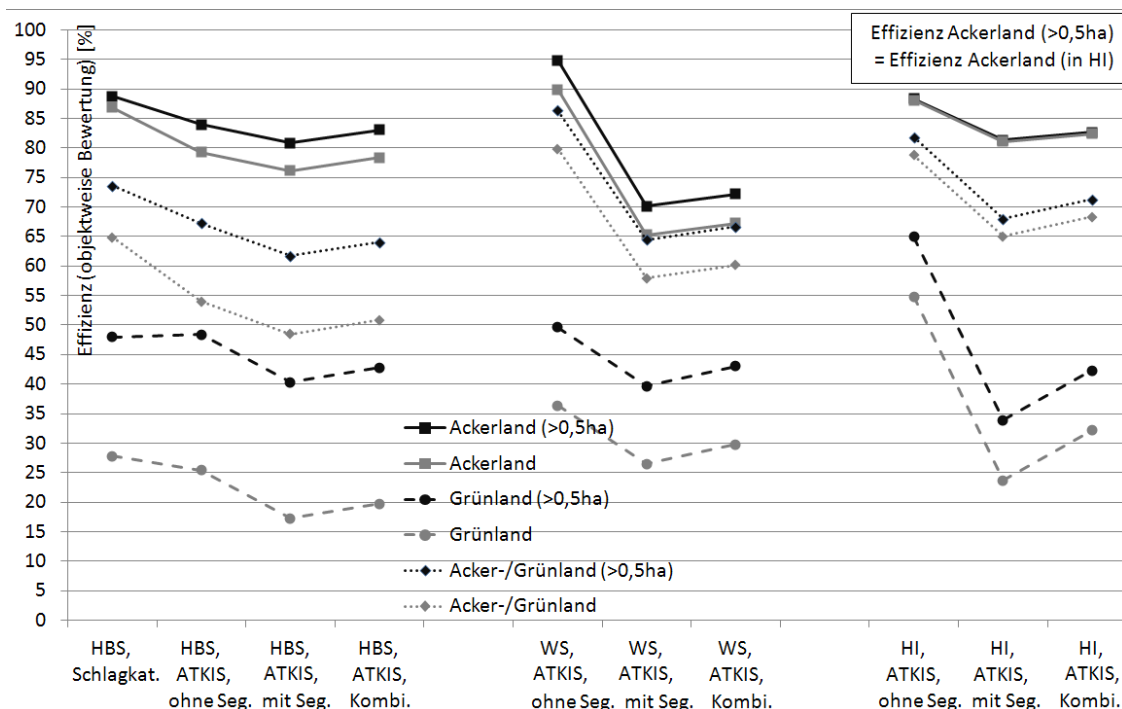


Abbildung 64: Evaluierungsergebnisse der Verifikation von GIS-Ackerland- und Grünlandobjekten sowie der Vereinigungsmenge aller Acker- und Grünlandobjekte des zweiten Analyseschrittes bzgl. der Effizienz (objektweise Betrachtung) unter Betrachtung verschiedener Flächengrößen.

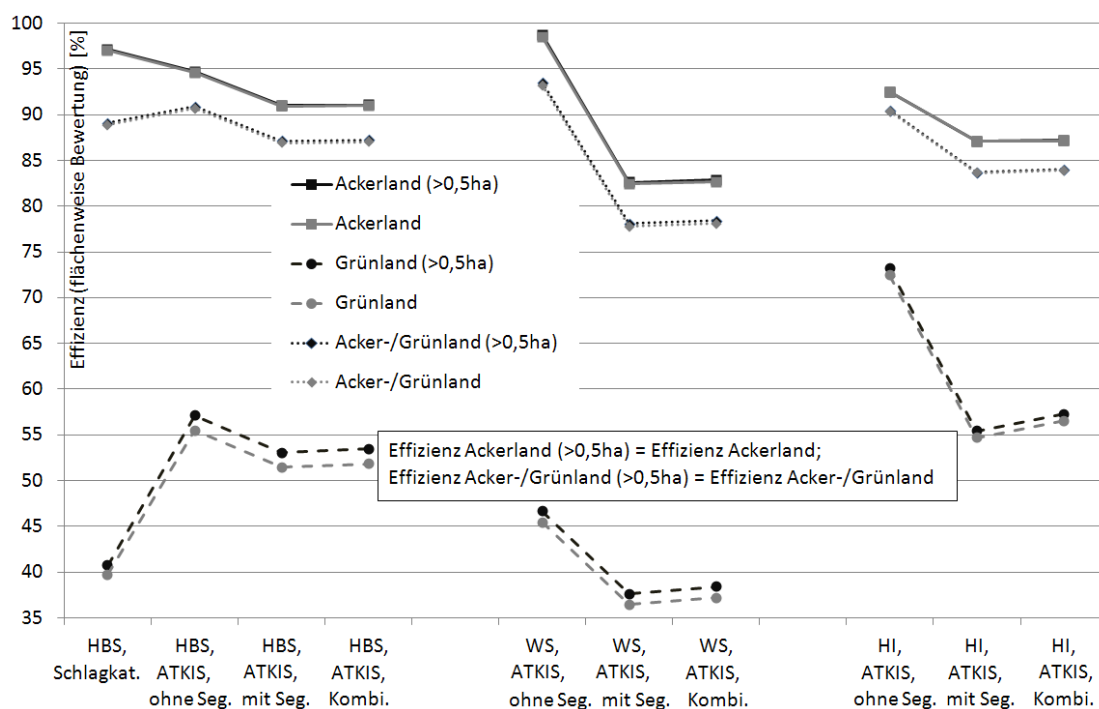


Abbildung 65: Evaluierungsergebnisse der Verifikation von GIS-Ackerland- und Grünlandobjekten sowie der Vereinigungsmenge aller Acker- und Grünlandobjekte des zweiten Analyseschrittes bzgl. der Effizienz (flächenweise Betrachtung) unter Betrachtung verschiedener Flächengrößen.

In den Abbildung 64 und Abbildung 65 sind neben der *Effizienz*, die sich auf alle GIS-Ackerland- und Grünlandobjekte bezieht, auch die eingetragen, bei denen die kleinen GIS-Objekte, also GIS-Objekt mit einer Fläche unter 0,5 ha, nicht berücksichtigt werden. GIS-Objekte kleiner als 0,5 ha sind wie in Abschnitt 5.3 beschrieben auf Grund ihrer Größe nicht klassifizierbar und werden daher grundsätzlich vom System

verworfen. Eine Verbesserung der *Effizienz*, bezogen auf die Bewertung der Anzahl der Objekte, ist die logische Folge, während sich die *Effizienz*, bezogen auf die Fläche, natürlich kaum ändert.

In Abbildung 66 sind die Evaluierungsergebnisse farbcodiert exemplarisch für die Testszene Halberstadt unter Verwendung von ATKIS (links: ohne Segmentierung, rechts: mit Segmentierung) dargestellt.

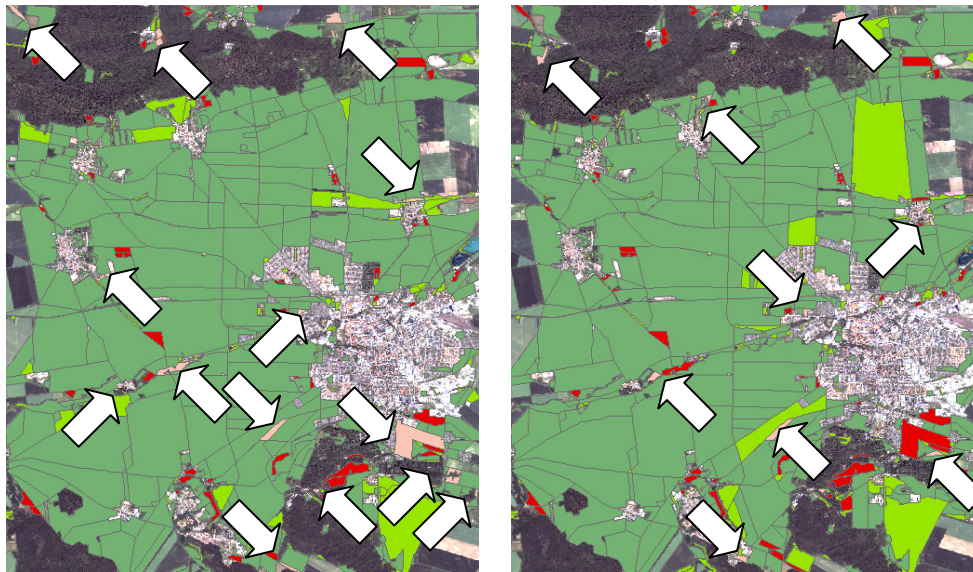


Abbildung 66: Evaluationsergebnisse der Verifikation von ATKIS der Szene HBS (links: ohne Segmentierung, rechts: mit Segmentierung) (TP: dunkelgrün, FN: hellgrün, FP: hellrot (durch Pfeile gekennzeichnet), TN: dunkelrot).

5.7.3. Bewertung

Das Hauptziel der Verifikation, das Erreichen einer *TG a posteriori* von 95%, bezogen auf die Fläche, konnte in allen Testdatensätzen erreicht werden. Die *TG a posteriori* von 95%, bezogen auf die Anzahl der Objekte, konnte nur in einem Fall, in der Szene Weiterstadt, unter Anwendung der Segmentierung sehr knapp nicht erreicht werden. Die Gründe hierfür wurden diskutiert. Das Ergebnis ist insgesamt zufriedenstellend.

Bezüglich der *Effizienz* ist anzumerken, dass der Vorteil der objektbasierten Klassifikation gegenüber einer pixelbasierten Klassifikation durch die Verwendung zusätzlicher Merkmale wie struktureller- oder geometrischer Merkmale (Abschnitt 2.1.1.3 und Abschnitt 4.2), bei sehr kleinen GIS-Objekten nicht genutzt werden kann. Im Gegenteil, da die zu klassifizierenden Flächen so klein sind, dass sie auf Grund ihrer Größe nicht klassifizierbar sind, werden diese der Zurückweisungsklasse zugeordnet. Da diese GIS-Objekte generell vom menschlichen Bearbeiter manuell kontrolliert werden müssen, sinkt die *Effizienz* des Systems. Dieser Einfluss verstärkt sich bei der Verwendung der Segmentierung. Ein Segment kann maximal so groß wie das GIS-Objekt sein. Liegt ein sehr kleines GIS-Objekt vor, sinkt die Größe der Fläche des zu klassifizierenden Segmentes unnötig. Ähnliche Schwierigkeiten treten natürlich auch bei GIS-Objekten auf, deren Schläge und damit auch deren Segmente so klein sind, dass sie nicht klassifizierbar sind (Abbildung 59). Die Schwierigkeiten bei der Klassifikation kleiner Segmente beruhen darauf, dass einige in dieser Arbeit verwendeten Merkmale eine gewisse Dimension des Segmentes voraussetzen. Das trifft besonders auf die strukturellen Merkmale zu. Die *Effizienz* der Kombinationsszenarien konnte gegenüber den Testszenarien mit Anwendung der Segmentierung gesteigert werden. Eine weitere Verbesserung des Ergebnisses ist jedoch wünschenswert. Zwei Beispiele von GIS-Grünlandobjekten mit einer Fläche kleiner als 1 ha sind in Abbildung 67 zu sehen. Diese sehr kleinen GIS-Objekte befinden sich oft in der Nähe von linienhaften GIS-Objekten wie Straßen oder Flüssen, wobei die Formen der GIS-Objekte selber stark variieren können. Bei der Klassifikation kleiner GIS-Objekte sollte daher auf Merkmale verzichtet werden, die eine Mindestflächengröße benötigen. Stattdessen könnte Nachbarschaftswissen eingeführt werden, welches u.a. direkt durch das GIS gegeben ist. Liegt ein sehr kleines GIS-Objekt zwischen Straßen oder an Straßen, Eisenbahnwegen bzw. Flüssen oder ist es komplett von Siedlung oder Industrie umgeben, handelt es sich erfahrungsgemäß um ein GIS-Grünland- und nicht um ein Ackerlandobjekt. Ob Nachbarschaftswissen bei großen GIS-Ackerland- und Grünlandobjekten das Ergebnis der Klassifikation steigern kann, ist allerdings zweifelhaft, da große GIS-Ackerland- oder Grünlandobjekte sowohl Bahnlinien als auch Straßen, Siedlungen oder Industriegebiete benachbart sind. Außerdem liefert die objektbasierte Klassifikation für große GIS-

Ackerland- oder Grünlandobjekte, besonders unter Verwendung von spektralen und strukturellen Merkmalen, sehr gute Ergebnisse.



Abbildung 67: Zwei Beispiele für GIS-Grünlandobjekte unter 1 ha Flächengröße (Testszene: HI).

5.8. Evaluierung des gesamten Algorithmus

In diesem Abschnitt erfolgt die Evaluation des gesamten Algorithmus, d.h. der Kombination des pixelbasierten Klassifikationsansatzes mit der Überführung der Ergebnisse auf GIS-Objekte zur Trennung der gemeinsamen Klasse Acker- und Grünland von allen anderen Klassen und der objektbasierten Klassifikation zur Trennung von GIS-Ackerland- und Grünlandobjekten mit oder ohne Anwendung der Segmentierung. Für die Evaluierung in diesem Abschnitt werden dieselben Testszenen benutzt, die bereits in den Abschnitten 5.4 und 5.7 verwendet wurden; es ergeben sich die selben Testszenarien wie in Tabelle 30 gezeigt. Für die Experimente wird der Parameter 2 (Schwellwert s_q) für die Überführung der pixelbasierten Klassifikation auf ein GIS-Objekt unter Berücksichtigung der Untersuchungen in Abschnitt 5.4 auf 30% festgelegt. In die Analyse des Verifikationsprozesses werden auch die GIS-Objekte einbezogen, die auf Grund der Größe vom zweiten Analyseschritt, der objektbasierten Klassifikation zur Trennung von Acker- und Grünland, gar nicht bewertet werden können und daher vom System verworfen werden müssen. Dies hat eine nicht geringfügig Steigerung der *False Negatives* zur Folge.

5.8.1. Evaluierung der Verifikationsergebnisse

Die Evaluierungsergebnisse der Verifikation der *TG a priori* und *TG a posteriori* aller Testszenarien sind in Abbildung 68 (objektweise Betrachtung) und Abbildung 69 (flächenweise Betrachtung) zusammengefasst. Die *TG a priori* weicht natürlich von Abbildung 40 (Abschnitt 5.5, S.72) und Abbildung 60 (Abschnitt 5.7, S.89) ab, da nun sowohl Verwechslungen von GIS-Ackerland- und Grünlandobjekten mit anderen GIS-Objekten als auch zwischen Acker- und Grünland zugelassen sind und detektiert werden müssen. Die *TG a priori* (objektweise Bewertung) liegt für die Klasse Grünland in allen Tests unter 95%, stellenweise sogar erheblich (Weiterstadt: 52,3%, Halberstadt (Schlagkataster): 61,7% und Halberstadt (ATKIS): 66,4%). Die geforderte *TG a posteriori* von 95% (objektweise Betrachtung) lag in der Testszene Hildesheim bereits vor (98,1%) und wird auch in Weiterstadt nur geringfügig überschritten (95,6%). Die *TG a priori* von Ackerland, bezogen auf die Fläche, liegt immer über den geforderten 95% und ist immer höher als die *TG a priori* dieser Klasse, bezogen auf eine objektweise Betrachtung. Auch damit wird deutlich, dass die fehlerhaften GIS-Ackerlandobjekte eher kleine GIS-Objekte sind. Ein einheitlicher Trend bezüglich der fehlerhaften GIS-Grünlandobjekte kann nicht beobachtet werden. Während es sich in der Testszene Weiterstadt und Hildesheim um flächenmäßig große fehlerhafte GIS-Objekte handelt, kann dies mit der Szene Halberstadt nicht bestätigt werden. Die geforderte *TG a posteriori* von 95% der GIS-Ackerlandobjekte, bezogen auf die Fläche in allen Testszenen wurde bereits a priori erreicht; die *TG a priori* der GIS-Grünlandobjekte liegt immer unter diesem Wert. In der Szene Halberstadt liegt die *TG a priori* (flächenweise Betrachtung) der Klasse Grünland sogar bei nur 56,9% (Schlagkataster) bzw. 87,8% (ATKIS). Positiv ist festzustellen, dass die geforderte *TG a posteriori* von 95%, bezogen auf die Anzahl der Objekte und die Fläche, in allen Tests erreicht wird. Im Vergleich zum vorhergehenden Abschnitt ist diesmal in allen Testszenen zu beobachten, dass die *TG a posteriori* unter Verwendung der Segmentierung gegenüber der Verifikation ohne Segmentierung steigt. Bei einer Kombination der Ergebnisse ohne und mit Segmentierung fällt die *TG a posteriori* gegenüber der Verifikation unter Anwendung der Segmentierung nicht ab.

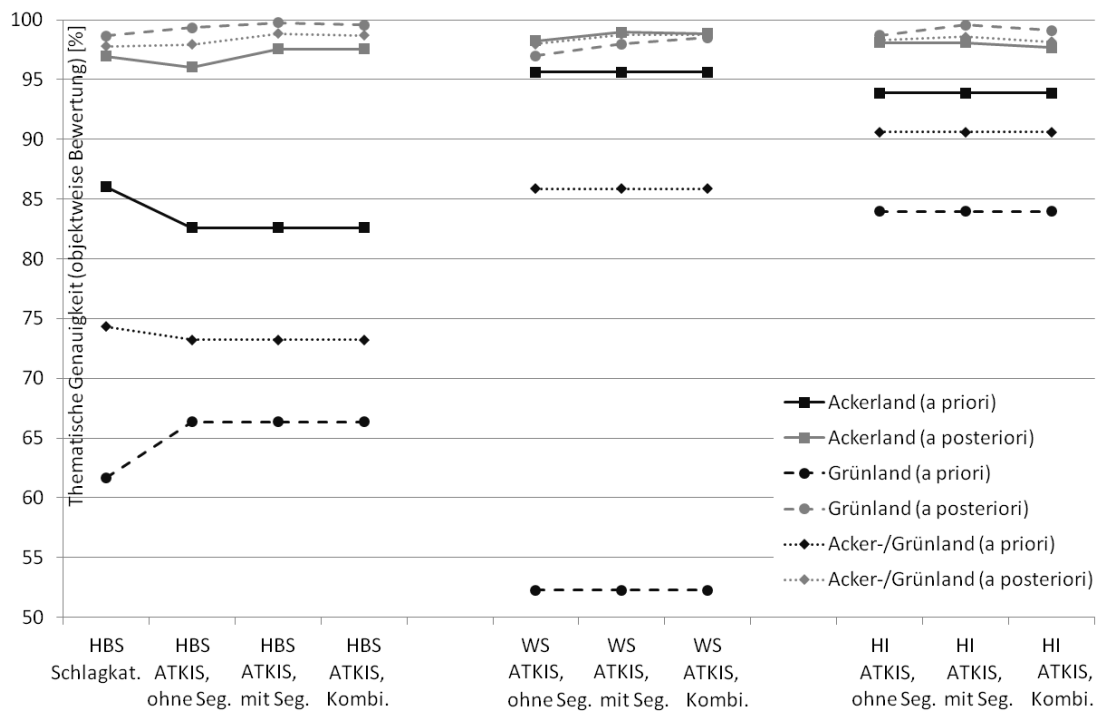


Abbildung 68: Evaluierungsergebnisse d. Verifikation der Klasse Ackerland, Grünland und der Vereinigungsmenge aller Acker-u. Grünlandobjekte bzgl. der *TG a priori* und *TG a posteriori* (objektweise Bewertung) des gesamten Algorithmus.

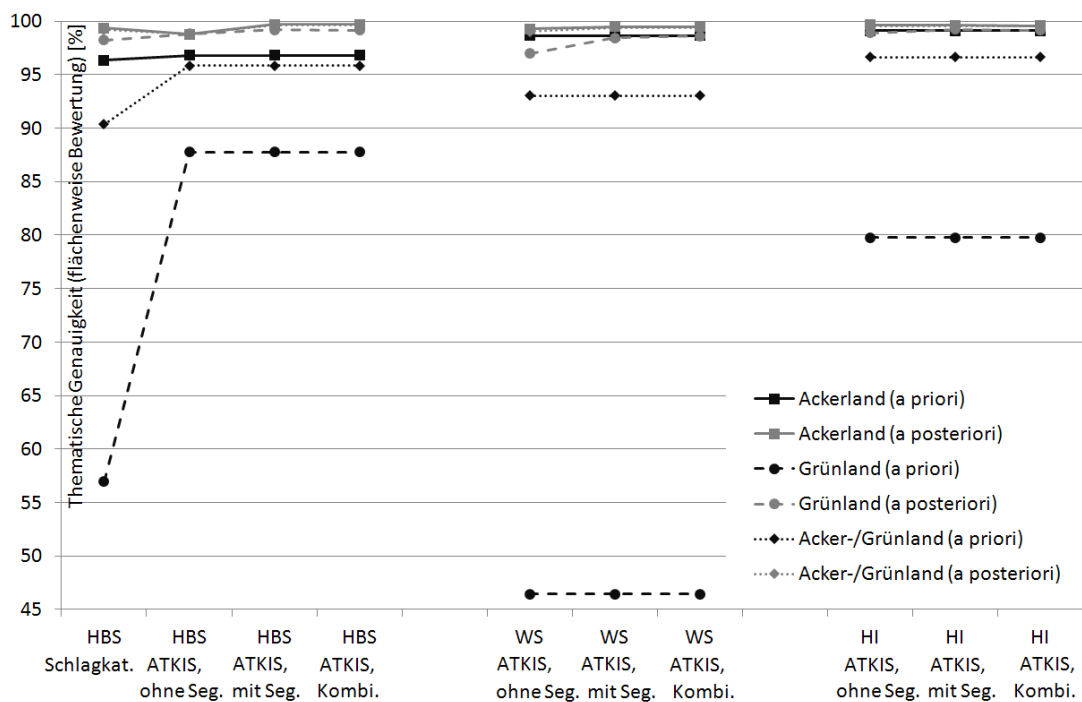


Abbildung 69: Evaluierungsergebnisse der Verifikation der Klasse Ackerland, Grünland und der Vereinigungsmenge aller Ackerland- und Grünlandobjekte bzgl. der *TG a priori* und *TG a posteriori* (flächenweise Bewertung).

In Abbildung 70 sind die *error detection rates* für alle Testszenarien sowohl bezüglich der Anzahl der Objekte als auch der Fläche dargestellt. Oft ist die *error detection rate* der Klasse Ackerland, bezogen auf die Anzahl der Objekte, höher als bezogen auf die Fläche, was ebenfalls für eher kleine fehlerhafte GIS-Objekte spricht. Bei Grünland ist es umgekehrt. Dies deutet daher auf eher größere fehlerhafte GIS-Objekte hin. Auch eine andere zuvor festgestellte Beobachtung, nämlich dass, wenn die *TG a priori* einer Klasse sehr hoch ist, die *error detection rate* dieser Szenarien geringer ist, kann bestätigt werden. Dies trifft in Abbildung 70 auf die

TG a priori der Klasse Ackerland in den Szenen Weiterstadt und Hildesheim zu. Bei diesen Beispielen fällt auch die *error detection rate* der Kombinationsszenarien deutlich ab, während sie in der Szene Halberstadt stagniert. Eine *error detection rate* von über 80% der Vereinigungsmenge von Ackerland- und Grünlandobjekten in allen Testszenarien, bezogen auf die Anzahl der Objekte, ist ein sehr zufriedenstellendes Ergebnis.

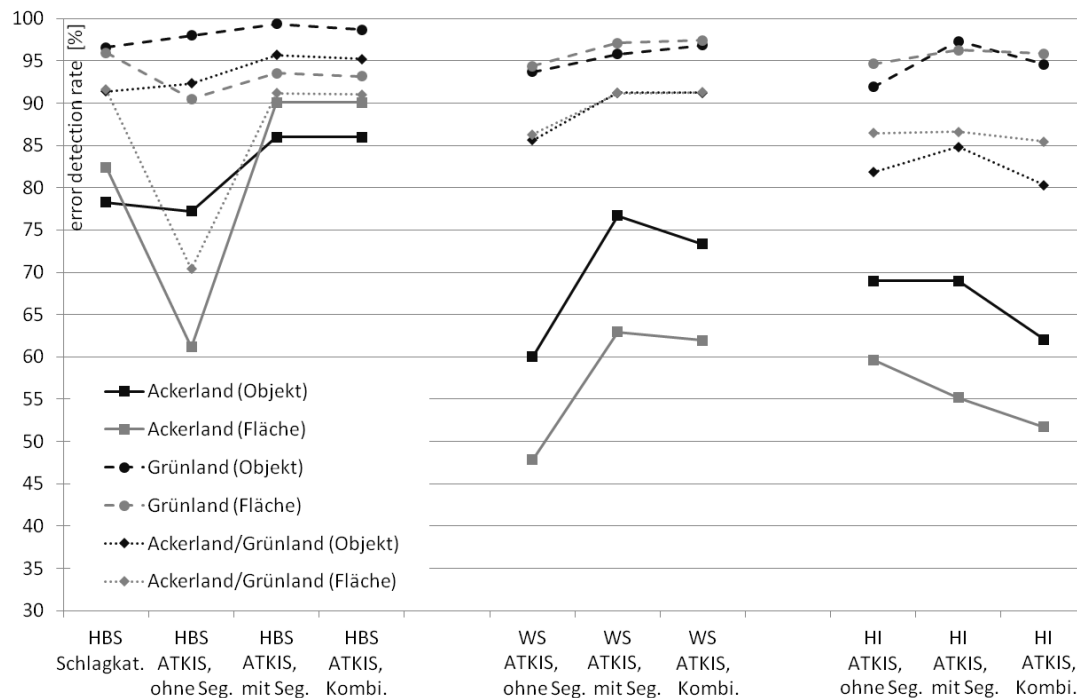


Abbildung 70: Evaluierungsergebnisse der Verifikation der Klasse Ackerland, Grünland und der Vereinigungsmenge der Acker- und Grünlandobjekte bzgl. der *error detection rate*.

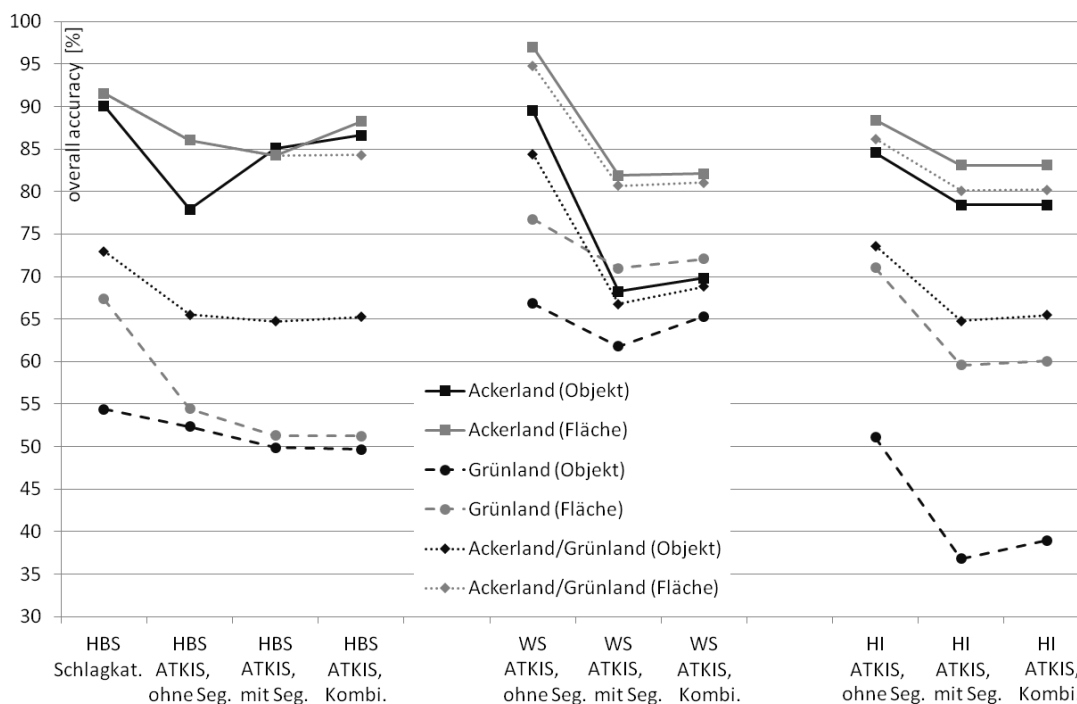


Abbildung 71: Evaluierungsergebnisse der Verifikation der Klasse Ackerland, Grünland und der Vereinigungsmenge der Acker- und Grünlandobjekte bzgl. der *overall accuracy*.

In Abbildung 71 sind die *overall accuracies* aller Testszenarien, bezogen auf die Anzahl der Objekte und der Fläche, graphisch dargestellt. Die *overall accuracies* der Klasse Grünland sind wie bei vorherigen Versuchen

i.d.R. geringer als die der Klasse Ackerland. Weiterhin fällt auf, dass die *overall accuracy*, bezogen auf die Anzahl der Objekte, immer kleiner als bezogen auf die Flächengröße ist, was ein weiterer Hinweis auf kleine fehlerhafte GIS-Objekte ist. Während die *error detection rate* der Vereinigungsmenge aus Acker- und Grünland bei den Kombinationsszenarien abfällt, steigt die *overall accuracy* an bzw. stagniert. Die *overall accuracy* der Vereinigungsmenge aus Acker- und Grünland, bezogen auf die Anzahl der Objekte, liegt immer über 65%, was ein sehr zufriedenstellendes Ergebnis darstellt. Nur 35% der Objekte wurden falsch zugeordnet.

In Abbildung 72 sind die Verifikationsergebnisse unter dem Aspekt der *Effizienz* dargestellt. Es fällt zunächst auf, dass die *Effizienz*, bezogen auf die Fläche (hellgrau), immer höher ist als die *Effizienz*, bezogen auf die Anzahl der Objekte (schwarz). Je weiter die jeweiligen Linien voneinander entfernt liegen, desto kleiner sind die manuell zu kontrollierenden GIS-Objekte. Während die *TG a posteriori* mit Anwendung der Segmentierung nur geringfügig gesteigert werden, fällt die *Effizienz* teils erheblich ab (einzige Ausnahme ist die *Effizienz* der Klasse Ackerland in der Testszene Halberstadt). Des Weiteren fällt, wie auch bei vorangegangenen Versuchen auf, dass die *Effizienz* der Klasse Ackerland immer höher als die *Effizienz* der Klasse Grünland ist. Die *Effizienz* der Klasse Grünland liegt unter Betrachtung der Fläche zwischen 20,5% und 53,0% und unter Betrachtung der Anzahl der Objekte nur zwischen 18,1% und 37,7%. Betrachtet man die Vereinigungsmenge der Ackerland- und Grünlandobjekte, steigt die *Effizienz* auf Grund der hohen Anzahl von GIS-Ackerlandobjekten. Die *Effizienz*, bezogen auf die Fläche der Vereinigungsmenge der Ackerland- und Grünlandobjekte, liegt zwischen 77,7% und 89,7% und bezüglich der Anzahl der Objekte zwischen 40,3% und 67,6%. Soll eine hohe *TG a posteriori* erreicht werden, z.B. 95%, und ist gleichzeitig die *TG a priori* geringer, kann ein hoher Anteil manueller Arbeit nicht vermieden werden. Ist die *TG a priori* höher, ist daher auch die *Effizienz* höher (Szene Hildesheim). Trotzdem erreicht der menschliche Bearbeiter immer eine Zeitersparnis unter Benutzung des Verifikationssystems, verglichen mit einer Verifikation ohne System.

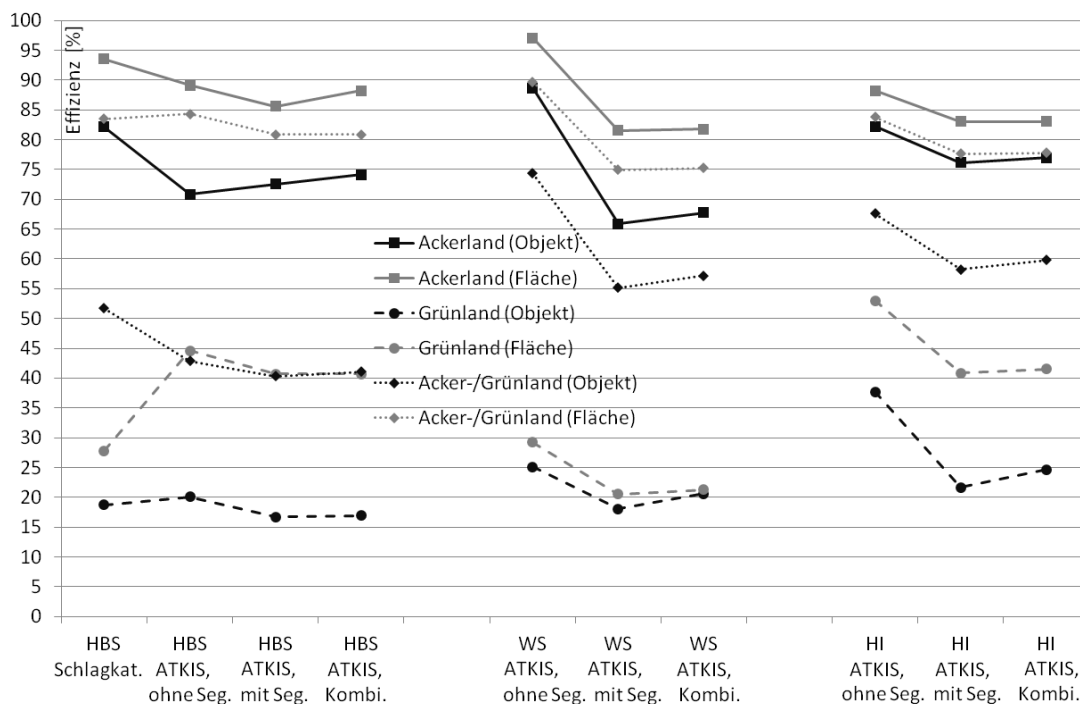


Abbildung 72: Evaluierungsergebnisse d. Verifikation der Klasse Ackerland, Grünland und der Vereinigungsmenge der Acker- und Grünlandobjekte bzgl. der *Effizienz* (objekt- und flächenweise Bewertung) des gesamten Algorithmus.

In Abbildung 73 erfolgt die Gegenüberstellung des Erfolgs der Verifikation, wenn zum einen alle GIS-Grünlandobjekte und zum anderen nur GIS-Grünlandobjekte mit einer Mindestgröße von 0,5 ha bzw. 1 ha betrachtet werden. In dieser Abbildung werden nur GIS-Grünlandobjekte in Bezug auf die Anzahl der Objekte dargestellt, da zuvor festgestellt wurde, dass die *Effizienz* besonders in Bezug auf die Anzahl der Objekte der Klasse Grünland eher gering ist. In Abbildung 73 wird deutlich, dass die *Effizienz* (schwarz) steigt und die *TG a posteriori* (mittlerer Grauton) konstant oder annähernd konstant bleibt, wenn kleinere Grünlandobjekte nicht betrachtet werden (davon abgesehen, dass die Mindestkartierfläche dieser Objektklasse im ATKIS Objektartenkatalog mit 1 ha festgeschrieben ist). Eine annähernd konstante *TG a*

posteriori bedeutet auch eine annähernd konstante *error detection rate*, auf deren explizite Darstellung verzichtet wird. Die Steigerung der *Effizienz*, wenn kleine GIS-Objekte nicht betrachtet werden, ist nicht nur auf die höhere *TG a priori* zurückzuführen, sondern auf die wesentlich geringere *False Negative Rate*, was auch an der Steigerung der *overall accuracy* zu sehen ist (Ausnahme ist die Szene Halberstadt, wo die *overall accuracy* leicht abfällt). Die *False Negative Rate* ist im Durchschnitt über alle Testszenen ca. 25% geringer bei der Betrachtung ohne GIS-Objekte unter 0,5 ha/1 ha gegenüber der Auswertung aller GIS-Objekte.

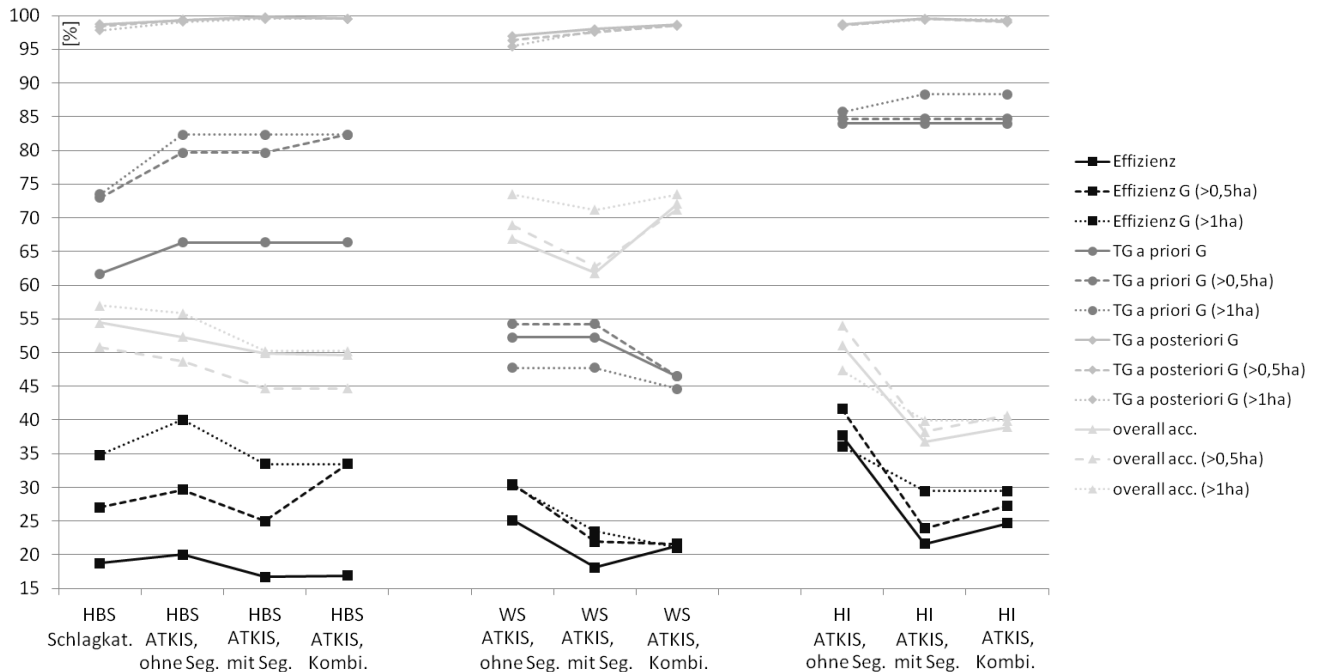


Abbildung 73: Evaluierungsergebnisse der Verifikation der Klasse Grünland (objektweise Bewertung) bzgl. Effizienz, TG a priori, TG a posteriori und overall accuracy unter Betrachtung verschiedener Flächengrößen.

Auch diesmal werden zur Vermittlung eines besseren Eindrucks die Evaluierungsergebnisse graphisch dargestellt (Abbildung 74). Dabei wird auch wie bisher die Testszene Halberstadt verwendet, die bei den Tests in diesem Abschnitt das schlechteste Ergebnis erzielte.

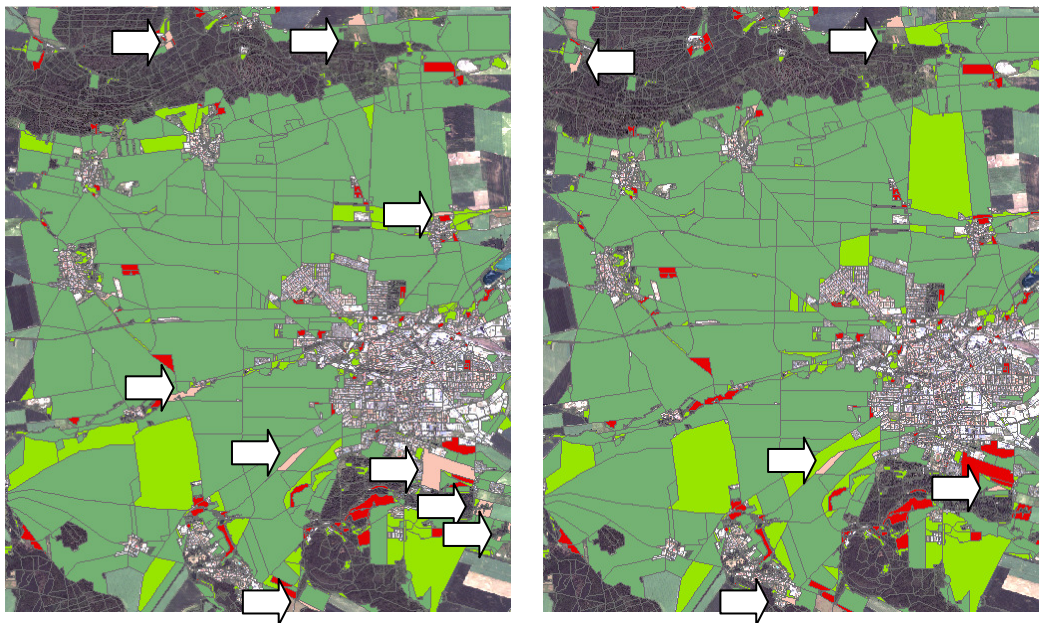


Abbildung 74: Evaluierungsergebnisse der Verifikation von ATKIS der Szene HBS (links: ohne Segmentierung, rechts: mit Segmentierung) (TP: dunkelgrün, FN: hellgrün, FP: hellrot (durch Pfeile gekennzeichnet), TN: dunkelrot).

5.8.2. Bewertung

Der Algorithmus zur Verifikation von GIS-Ackerland- und Grünlandobjekten konnte in allen Tests die gestellten Anforderungen erreichen; die *TG a posteriori* sowohl der Klasse Ackerland als auch der Klasse Grünland konnten die in (BKG, 2009) gestellten Anforderungen erfüllen (*TG* von 95%, bezogen auf die Anzahl der Objekte und auf die Fläche). Unter Verwendung der Segmentierung konnte die *TG a posteriori* sowohl unter Betrachtung der Anzahl der Objekte als auch unter Betrachtung der Fläche noch weiter gesteigert werden; gleichzeitig sank die *Effizienz* jedoch deutlich. Bei den Kombinationsszenarien konnte die *Effizienz* wieder leicht gesteigert werden. Der menschliche Bearbeiter muss also abwägen zwischen einer sehr hohen *TG a posteriori* und dem Aufwand der manuellen Kontrolle der vom System abgelehnten Objekte. Wenn viele Fehler in einem GIS vorhanden sind, ist die *Effizienz* ebenfalls gering, da korrekterweise diese Fehler gefunden werden. Dies ist bei den evaluierten Datensätzen bei der Klasse Grünland besonders häufig der Fall. Ebenfalls wird die *Effizienz* durch besonders viele kleine GIS-Objekte bzw. Schläge in den GIS-Objekten gedrückt. Unterschreiten die GIS-Objekte bzw. Segmente eine Mindestgröße, die für die objektbasierte Klassifikation notwendig ist, sind diese nicht klassifizierbar, sie werden vom System abgelehnt und müssen somit manuell kontrolliert werden. Dabei ist sehr auffällig, dass besonders viele GIS-Grünlandobjekte existieren, die kleiner als die im GIS-Katalog festgelegten Mindestkartierflächen sind. Fazit: Es können 50-60% der Arbeit gespart werden und dabei die *TG a posteriori* auf über 95% steigern.

5.9. Zusammenfassung

In diesem Kapitel wurde die Methode zur Verifikation von GIS-Ackerland- und Grünlandobjekten evaluiert. Zunächst wurde der erste Analyseschritt, der die gemeinsame Klasse Acker-/Grünland von anderen Klassen trennt und auf den Arbeiten von Gimel'farb (1996) und Busch et al. (2005) beruht, evaluiert. Es ist wichtig, dass dieser Analyseschritt sehr zuverlässige Ergebnisse mit einer hohen *TG a posteriori* liefert, da Fehler, die bei diesem Schritt unentdeckt bleiben, vom zweiten Analyseschritt nicht aufgedeckt werden können. Das Ergebnis des gesamten Algorithmus kann daher niemals besser als das des ersten Analyseschrittes sein. Die gestellten Anforderungen (BKG, 2009) konnten erfüllt werden, da eine *TG a posteriori* von mehr als 95% in allen Tests erreicht wird. Die Parameter für die pixelbasierte Klassifikation werden dabei automatisch trainiert (Becker et al., 2009), die Parameter der Überführung der pixelbasierten Klassifikation auf GIS-Objekte gehen auf den GIS-Objektartenkatalog und Erfahrungswerte zurück. Den größten Einfluss hat der Parameter s_q , der mit Hilfe von Wissen über die Erscheinungsformen von Acker-/Grünland in einer Szene eingestellt werden kann. Eine Testreihe hat die Erfahrungswerte von 30% - 40% bestätigt.

Der zweite Analyseschritt ist eine objektbasierte Klassifikation der Ackerland- und Grünlandobjekte, die nicht vom ersten Analyseschritt abgewiesen wurden. Ziel der objektbasierten Klassifikation ist, zwischen Acker- und Grünland zu unterscheiden. Dieser Analyseschritt ist neuartig und daher intensiver getestet. Zunächst wurden die Untersuchungen auf ein Schlagkataster durchgeführt, um den Kern des zweiten Analyseschrittes, die objektbasierte Klassifikation mittels SVM und die dafür benutzten Merkmale, detailliert zu evaluieren. Alle Parameter dieses Teils der objektbasierten Klassifikation werden entweder automatisch gelernt oder werden einmalig durch empirische Tests auf Werte festgelegt, die dann in allen Tests gleichermaßen verwendet wurden. In der ersten Untersuchung wurde zunächst der Einfluss der unterschiedlichen Merkmalsgruppen auf den Erfolg der Klassifikation und Verifikation evaluiert. Die jeweils besten Ergebnisse werden erreicht, wenn spektrale und strukturelle Merkmale genutzt werden. Das beste Ergebnis wird bei der Verwendung von Merkmalen aus allen Merkmalsgruppen erzielt. Die geforderte *TG a posteriori* von 95% konnte in allen Kombinationen erreicht werden, dabei schwankt besonders die *Effizienz* der Klasse Grünland.

In der zweiten Untersuchung bezüglich der objektbasierten Klassifikation wurde die Abhängigkeit des Erfolges der Klassifikation und Verifikation vom Aufnahmezeitpunkt des Bildes betrachtet. Eine Abhängigkeit des Erfolgs für die Klassifikation und Verifikation von den Zeitpunkten wurde festgestellt. Bei dieser Testreihe haben sich unterschiedliche Zeitpunkte für Klassifikation und Verifikation als besonders geeignet herausgestellt. Insgesamt kann jedoch keine abschließende Bewertung vorgenommen werden, da der GIS-Datensatz zum einen nur wenige Ackerland- und Grünlandobjekte umfasste und zum anderen es Kalibrierungsprobleme beim NIR-Kanal gab. Weitere Tests der zeitlichen Abhängigkeit des objektbasierten Klassifikationsverfahrens zur Trennung von Acker- und Grünland mit einem größeren GIS-Datensatz und multispektralen Bildern sind ratsam.

Anschließend wurde der zweite Analyseschritt auf GIS-Daten, die dem Inhalt und Detailierungsgrad einer topographischen Karte mittleren Maßstabs entsprechen, angewendet, wobei auf GIS-Objekte fokussiert wurde, die entweder der Klasse Acker- oder Grünland angehörten und keine Fehler enthielten, die durch den erste Analyseschritt aufzudecken sind. Die Parameter für die Segmentierung wurden empirisch bestimmt, dabei wurde zwischen den verschiedenen Tests nur eine kleine Anpassung vorgenommen. Das Hauptziel der Verifikation, das Erreichen einer *TG a posteriori* von 95%, bezogen auf die Fläche, wurde in allen Testszenen erreicht. Die Bedingung einer *TG a posteriori* von 95%, bezogen auf die Anzahl der Objekte, wurde nur in einem Fall knapp nicht erfüllt. Der Grund hierfür lag an teilweise falsch klassifizierten Segmenten innerhalb von GIS-Objekten. Die *Effizienz* sinkt bei der Benutzung der Segmentierung des zweiten Analyseschrittes deutlich gegenüber einer Klassifikation ohne Segmentierung ab, während die *TG a posteriori* leicht steigt. Mit Hilfe der Kombination der Ergebnisse ohne und mit Segmentierung (Kombinationsszenarien) wird die *Effizienz* wieder verbessert. Bei den Kombinationsszenarien konnte immer eine *TG a posteriori* von über 95% erreicht werden, was bei der Verwendung der Segmentierung auf alle GIS-Objekte in einem Fall nicht möglich war.

Abschließend wurde der gesamte Algorithmus, bestehend aus der Kombination beider Analyseschritte, auf reale GIS-Daten angewendet. Der Algorithmus erreichte in allen Tests die geforderte *TG a posteriori* von 95%, bezogen auf die Anzahl der Objekte und die Fläche, sowohl für die Klasse Ackerland als auch für die Klasse Grünland. Unter Verwendung der Segmentierung wird die *TG a posteriori* sowohl unter Betrachtung der Anzahl der Objekte als auch unter Betrachtung der Fläche gesteigert; zeitgleich sinkt jedoch die *Effizienz*. Bei einer kombinierten Auswertung (Kombinationsszenarien) steigt die *Effizienz* wieder an, erreicht aber nicht das Niveau wie bei einer Auswertung komplett ohne Segmentierung. Der menschliche Bearbeiter muss also abwägen zwischen einer sehr hohen *TG a posteriori* und dem Aufwand der manuellen Kontrolle der vom System abgelehnten Objekte. Ebenfalls sinkt die *Effizienz* durch besonders viele kleine GIS-Objekte, da für die objektbasierte Klassifikation eine Mindestgröße notwendig ist und GIS-Objekte, die auf Grund ihrer Größe nicht klassifizierbar sind, vom System abgelehnt und manuell kontrolliert werden müssen. Dabei ist besonders auffällig, dass viele GIS-Grünlandobjekte existieren, die die im GIS-Katalog festgelegte Mindestkartierfläche unterschreiten.

6. Folgerungen und Ausblick

Das Ziel der Arbeit, die Entwicklung einer Methode für die semi-automatische Verifikation von GIS-Objekten der Objektarten Acker- und Grünland unter Verwendung orthorektifizierter monotemporaler Luft- und Satellitenbilder mit einer geometrischen Auflösung von 0,5 m bis 1m, wurde, wie in Kapitel 5 gezeigt, erreicht. Beurteilt wurde der Algorithmus nach den Kriterien für die Erfassung des DLM-DE, einem GIS, das den Inhalt und Detaillierungsgrad einer topographischen Karte mittleren Maßstabs entspricht (BKG, 2009). Die Kriterien in (BKG, 2009) lauten: eine *TG* von 95%, bezogen auf die Anzahl der Objekte und die Fläche, ist zu erreichen. In allen Testszenen wird die *TG a posteriori* gegenüber der *TG a priori* gesteigert, wobei für die *TG a posteriori* immer die geforderte Marke von 95% erreicht wird (Tabelle 32). Die dabei erzielte *Effizienz* schwankt zwischen 40% und 65%, bezogen auf die Anzahl aller Objekte der Klassen Acker- und Grünland. Verglichen zu bisher existierenden Verifikationssystemen konnten erstmals die getrennten Klassen Acker- und Grünland verifiziert werden; dem Benutzer entsteht eine Zeitersparnis von über 50% bei der Verwendung des vorgestellten Verifikationsansatzes.

Objektklasse	TG bezogen auf	TG a priori	TG a posteriori
Grünland	Anzahl der Objekte	67,0 %	98,8 %
	Fläche	69,9 %	98,7 %
Ackerland	Anzahl der Objekte	90,2 %	97,8 %
	Fläche	98,0 %	99,5 %
Vereinigungsmenge	Anzahl der Objekte	82,4 %	98,4 %
Acker-/Grünland	Fläche	94,7 %	99,4 %

Tabelle 32: Vergleich der durchschnittlichen *TG a priori* und *TG a posteriori* aller Testszenen.

Das vorgestellte Verifikationsverfahren gliedert sich in zwei Schritte. Im ersten Schritt wird zwischen der gemeinsamen Klasse *Acker-/Grünland* und weiteren Klassen wie *Siedlung*, *Industrie* und *Wald* unterschieden. Das Verfahren beruht auf dem Ansatz von Gimelfarb (1996). Ziel dieses Analyseschrittes ist es, GIS-Ackerland- und Grünlandobjekte zu detektieren, bei denen eine Teilfläche mit einer anderen Nutzung vorliegt, z.B. Gebäude einer Siedlungsfläche. Welche Teilflächen innerhalb eines GIS-Ackerland- und Grünlandobjektes nicht erlaubt sind, wird im Objektartenkatalog des GIS definiert. Teilflächen anderer Nutzung wurden von der pixelbasierten Klassifikation immer erfolgreich detektiert, aber bei der Überführung des pixelbasierten Ergebnisses auf ein GIS-Objekt nach Busch et al. (2005) kam es teilweise zu Informationsverlusten, auch wenn dies selten war. In Abschnitt 5.4.3 wurde bereits erörtert, dass eine einfache Anpassung der Parameter für diesen Schritt nicht ausreichend ist, da damit höchstens eine Anpassung an die Testszene erreicht werden kann. Eine Möglichkeit, diesem Problem zu begegnen, ist das Verwenden von Nachbarschaftsrelationen. Während z.B. korrekt detektierte Gebäude oft in unmittelbarer Nähe zu anderen Gebäuden liegen (Abbildung 42, S. 73), die ebenfalls detektiert wurden, werden um fälschlich detektierte Gebäude zu meist keine weiteren Gebäude detektiert (Abbildung 43). Diese Nachbarschaftsrelationen sollten bei der Überführung des pixelbasierten Klassifikationsergebnisses auf ein GIS-Objekt genutzt werden.

Alternativ oder zusätzlich können auch Höheninformationen aus einem normalisierten Digitalen Geländemodell (DGM_{norm}) verwendet werden, das durch Subtraktion eines Digitalen Oberflächenmodells (DOM) mit einem Digitalen Geländemodell (DGM) entsteht. DGM und DOM können aus Laserscanbefliegung oder Bildzuordnung von Stereobildpaaren abgeleitet werden. Mittels des DGM_{norm} kann dann verifiziert werden, ob eine z.B. als Industrie detektierte Fläche tatsächlich eine Industriefläche ist (Abbildung 42). Dies kann dadurch realisiert werden, indem überprüft wird, ob in dieser Fläche ein Gebäude vorliegt, es also eine Erhöhung im DGM_{norm} gibt.

Beim zweiten Analyseschritt, der Trennung von allen beim ersten Analyseschritt nicht abgelehnten GIS-Objekten in die Klassen Acker- und Grünland, ist auf Grund der Spezifikation von GIS, die den Inhalt und Detaillierungsgrad einer topographischen Karte mittleren Maßstabs entsprechen, die Segmentierung eines GIS-Objektes in verschiedene Schläge nötig. Der in dieser Arbeit verwendete Algorithmus arbeitet sehr zuverlässig und besitzt nur wenige Parameter, die leicht eingestellt werden können. Die Evaluierung des Ansatzes hat gezeigt, dass die *TG a posteriori* mittels Anwendung der Segmentierung gesteigert werden kann, die *Effizienz* jedoch sinkt, da viele Objekte der Zurückweisungsklasse zugeschrieben werden. Ein GIS-Objekt wird der Zurückweisungsklasse zugeschrieben, wenn mehr als die Hälfte der Fläche mit Segmenten

bedeckt ist, die auf Grund deren Größe nicht klassifiziert werden konnte oder die Fläche, die das GIS-Objekt von der Klasse Ackerland oder Grünland bedeckt, nicht die Mindestkartierfläche erreicht oder diese von beiden Klassen erreicht wird. Eine Zuweisung des GIS-Objektes zu der Zurückweisungsklasse tritt besonders häufig bei sehr kleinen GIS-Objekten auf. Eine Möglichkeit, dieses Problem nicht entstehen zu lassen, ist die Einführung eines Schwellwertes, ab wann eine Segmentierung eines GIS-Objektes vorgenommen werden soll. Wenn dieser Schwellwert unterschritten wird, wird auf eine Segmentierung der GIS-Objekte verzichtet. Tests haben gezeigt, dass dieser Ansatz zu einer Steigerung der *Effizienz* gegenüber dem Ansatz, bei der die Segmentierung für alle GIS-Objekte genutzt wird, führt. Eine vergleichbare *Effizienz*, die ohne die Anwendung der Segmentierung erzielt wurde, konnte jedoch nicht erreicht werden. Eine Möglichkeit, die *Effizienz* weiter zu steigern, ohne dass die *TG a posteriori* abnimmt, ist die Einführung zusätzlicher Informationen in die Klassifikation/ Verifikation. Die sehr kleinen GIS-Objekte, bei denen es hauptsächlich zu Schwierigkeiten bei der Klassifikation/Verifikation kommt, befinden sich oft in der Nähe von linienhaften GIS-Objekten wie Straßen oder Flüssen. Bei der Klassifikation kleiner GIS-Objekte könnten ebenfalls Nachbarschaftsrelationen eingeführt werden, die u.a. direkt durch das GIS gegeben sind. Liegt ein sehr kleines GIS-Objekt zwischen Straßen oder an Straßen, Eisenbahnwegen bzw. Flüssen oder ist es komplett von Siedlung oder Industrie umgeben, handelt es sich erfahrungsgemäß um GIS- Grünland- und nicht um ein Ackerlandobjekt. Ob Nachbarschaftsrelationen bei großen GIS-Ackerland- und Grünlandobjekten das Ergebnis der Klassifikation steigern können, ist allerdings anzuzweifeln, da große GIS-Ackerland- oder Grünlandobjekte sowohl von Bahnlinien als auch Straßen, Siedlungs- oder Industriegebiete benachbart sind. Außerdem liefert die objektbasierte Klassifikation für große GIS-Ackerland- oder Grünlandobjekte, besonders unter Verwendung von spektralen und strukturellen Merkmalen, sehr gute Ergebnisse.

Bei der Evaluierung wurde festgestellt, dass der Erfolg der objektbasierten Klassifikation des zweiten Analyseschrittes vom Zeitpunkt der Aufnahme der Bilddaten abhängig ist. Es ist zu prüfen, ob es Merkmale gibt, die eine Trennung von Acker- und Grünland innerhalb der Zeitpunkte ermöglicht, bei denen bisher keine befriedigenden Ergebnisse für Klassifikation bzw. Verifikation erreicht wurden. Die Hinzunahme dieser Merkmale in den Klassifikationsprozess kann dann eine zuverlässige Trennung zwischen Acker- und Grünland zu fast allen Zeitpunkten gewährleisten, was die Ergebnisse stabilisiert und zuverlässiger macht.

Abschließend ist zu sagen, dass diese Arbeit gezeigt hat, dass die Qualitätskontrolle der volkswirtschaftlich relevanten und komplexen GIS-Objekte der Klassen Acker- und Grünland automatisiert durchgeführt werden kann und damit einen wichtigen Beitrag zur Automatisierung der Qualitätskontrolle liefert. Forschungsbedarf besteht hinsichtlich der Überprüfung weiterer Qualitätsmaße wie der *geometrischen Genauigkeit* und der Überprüfung anderer Objektarten eines GIS, die in dieser Arbeit keine Berücksichtigung gefunden haben, um das Ziel der automatisierten Qualitätskontrolle von GIS-Daten zu verwirklichen, die dann eine zuverlässige Grundlage für die Entscheidungen des öffentlichen und privaten Lebens gewährleisten.

Literaturverzeichnis

- AdV, Arbeitsgemeinschaft der Vermessungsverwaltungen der Länder der Bundesrepublik Deutschland, 1997: ATKIS - Amtlich Topographisch-Kartographisches Informationssystem, Germany. In <http://www.atkis.de> (Datum des letzten Besuches: 28. Juni 2011).
- AdV, Arbeitsgemeinschaft der Vermessungsverwaltungen der Länder der Bundesrepublik Deutschland, 2011: Amtliches Topographisch-Kartographisches Informationssystem – ATKIS. In <http://www.adv-online.de/icc/extdeu/broker.jsp?uMen=ab9708a8-6975-7011-3bbc-251ec0023010> (Datum des letzten Besuches: 12. September 2011).
- Akca, D., Freeman, M., Sargent, I., Guen, A., 2010: Quality Assessment of 3D Building Data. In *Photogrammetric Record*, vol. 25, no. 132. S. 339–355.
- Albertz, J., 2009: Einführung in die Fernerkundung – Grundlagen der Interpretation von Luft- und Satellitenbildern. In *Wissenschaftliche Buchgesellschaft*, 4. aktualisierte Auflage, 254 Seiten.
- Alemán-Flores, Miguel, Álvarez-León, Luis, 2003: Texture classification through multiscale orientation histogram analysis. In *Proc. 4th international conference on Scale space methods in computer vision*, S. 479–493.
- Aplin, A., Atkinson, P.M., Curran, P.J., 1999: Fine Spatial Resolution Simulated Satellite Sensor Imagery for Land Cover Mapping in the United Kingdom. In *Remote Sensing of Environment*, vol. 68, S. 206–216.
- Alpin, P., Smith, G.M., 2008: Advances in Object-Based Image Classification. In *IntArchPhRS*, vol. 37, no. B7, S. 725–728.
- Amadsun, M., King, R., 1989: Textural Features Corresponding to Textural Properties. In *IEEE Transactions on Systems, Man, and Cybernetics*, vol. 19, no. 5, S. 1264–1274.
- Arkun, S.A., Dunk, I.J., Ranson, S.M., 2001: Hyperspectral remote sensing for vineyard management. In *Geospatial Information & Agriculture*, vol. 18, auf CD.
- Bach, F.R., Lanckriet, G.R.G., Jordan, M.I., 2004: Fast kernel learning using sequential minimal optimization. In *Report No. UCB/CSD-04-1307*, 2004.
- Balaguer, A., Ruiz, L. A., Hermosilla, T., Recio, J.A., 2010: Definition of a comprehensive set of texture semivariogram features and their evaluation for object-oriented image classification. In *Computers & Geoscience*, vol. 36, S. 231–240.
- Becker, C., Büschenfeld, T., Ostermann, 2009: Impacts of a Resolution Pyramid on Gibbs Random Field Classification. In *IntArchPhRS*, vol. 38, no. 1-4-7/WS, 4 Seiten, auf CD.
- Benz, U. C., Hofmann, P., Willhauck, G., Lingenfelder, I., Heynen, M., 2004: Multi-resolution, object-oriented fuzzy analysis of remote sensing data for GIS-ready information. In *ISPRS Journal of Photogrammetry and Remote Sensing*, vol. 58, S. 239–258.
- Bishop, C., 2006: Pattern Recognition and Machine Learning. In *Springer Verlag*, 738 Seiten.
- BKG (Bundesamt für Kartographie und Geodäsie), 2009: Leistungsbeschreibung zum Vorhaben “Aktualisierung des DLM-DE für das Stichjahr 2009”, In *Leistungsbeschreibung zum Vorhaben “Aktualisierung des DLM-DE für das Stichjahr 2009”*, Version 1.0 vom 12.02.2009, 27Seiten.
- Blascke, T., 2010: Object based image analysis for remote sensing. In *ISPRS Journal of Photogrammetry and Remote Sensing*, vol. 65, S. 2–16.
- Bogaert, J., Rousseau, R., Hecke, P.V., Impens, I., 2000: Alternative area-perimeter ratios for measurement of 2D shape compactness of habitats. In *Applied Mathematics and Computation*, vol. 111, no.1, S. 71–85.
- Bousquet, O. und Schölkopf, B., 2006: Comment. In *Statistical Science*, vol. 21, no. 3, S. 337–340.
- Brank, J., Grobelnik, M., Milic-Frayling, N., Mladenic, D., 2002: Feature Selection Using Linear Support Vector Machines. In *Technical Report, MSR-TR-2002-63*, Microsoft Research, Microsoft Corporation, 21 Seiten.
- Braun, A. C., Weidner, U., Hinz, S., 2010: Support Vector Machines for Vegetation Classification – A Revision. In *PFG*, vol. 4, S. 273–281.

- Bremer, M., Liebig, W., Prössler, S., 1992: Einrichtung des ATKIS-DLM 25/1 in Niedersachsen. In *Nachrichten der Niedersächsischen Vermessungs- und Katasterverwaltung*, no. 3, vol. 42, S. 134–157.
- Brügelmann, R., Förstner, W., 1992: Noise estimation for color edge extraction. In *Robust Computer Vision*, Wichmann-Verlag, S. 90–107.
- Burger, W., Burge, M. J., 2006: Digitale Bildverarbeitung: Eine Einführung mit Java und ImageJ. In *Springer Verlag*, 2. Auflage, 515 Seiten.
- Burges, C., 1998: A Tutorial on support Vector Machines for Pattern Recognition. In *Data Mining and Knowledge Discovery*, no.2, S. 121–167.
- Busch, A., Gerke, M., Grünreich, D., Heipke, C., Liedtke, C. E., Müller, S., 2004: Automated Verification of a Topographic Reference Dataset: System Design and Practical Results. In *IntArchPhRS*, vol. 35, no. B2, S. 735–740.
- Busch, A., Gerke, M., Grünreich, D., Heipke, C., Liedtke, C., Müller, S., 2005: Automatisierte Verifikation topographischer Geoinformation unter Nutzung optischer Fernerkundungsdaten. In *PFG*, vol. 2, S. 111–122.
- Canny, J.F., 1986: A computational approach to edge detection. In *IEEE TPAMI*, vol. 8, no. 6, S. 679–698.
- Carleer, A.P., Wolff, E., 2008: Change detection for updates of vector database through region-based classification of VHR satellite data. In *Proc. GEOBIA*, vol. 38, auf CD.
- Champion, N., Boldo, D., Pierrot-Deseilligny, M., Stamon, G., 2010: 2D building change detection from high resolution satellite imagery: A two-step hierarchical method based on 3D invariant primitives. In *Pattern Recognition Letters*, vol. 31, S. 1138–1147.
- Chanussot, J., Bas, P., Bombrun, L., 2005: Airborne Remote Sensing of Vineyards for the Detection of Dead Vine Trees. In *Proc. IGARSS*, S. 3090–3093.
- Cliff, A., Ord, J.K., 1981: Spatial Processes, Models and Applications. In *Pion Ltd*, 266 Seiten.
- Congalton, R., Green, K. 2009: Assessing the Accuracy of Remotely Sensed Data: Principles and Practices. In *CRC/Taylor & Francis*, 2. Auflage, 183 Seiten.
- Curran, P. J., 1988: The Semivariogram in Remote Sensing: An Introduction. In *Remote Sensing of Environment*, vol. 24, S. 493–507.
- Curran, P.J., 2001: Remote Sensing: Using the spatial domain. In *Environmental and Ecological Statistics*, vol. 8, S. 331–344.
- Delenne, C., Durrieu, S., Rabatel, G., Deshayes, M., Bailly, J.-S. Lelong, C., Couteron, P., 2007: Textural approaches for vineyard detection and characterization using very high spatial resolution remote sensing data. In *Int. J. of Remote Sensing*, vol. 29, no. 4, S. 1153–1167.
- Delenne, C., Rabatel, G., Deshayes, M., 2008: An Automatized Frequency Analysis For Vine Plot Detection And Delineation In Remote Sensing. In *IEEE Geoscience and Remote Sensing Letters*, vol.5, no.3, S. 341–345.
- De Wit, A. J. W. and Clevers, J. G. P. W., 2004: Efficiency and accuracy of per-field classification for operational crop mapping. In *Int. J. of Remote Sensing*, vol. 25, no. 20, S. 4091–4112.
- Duda, R.O., Hart, P.E., Stork, D. G., 2000: Pattern Classification. In *Wiley-Interscience Verlag*, 2.Auflage, 654 Seiten.
- Durrieu, M., Ruiz, L.A., Balaguer, A., 2005: Analysis of geospatial parameters for texture classification of satellite images. In *Proc. 25th EARSEL Symposium*, S. 11–18.
- Eckstein, W., 1996: Segmentation and texture analysis. In *IAPRSSIS*, vol. 31, no. B3, S. 165–175.
- EEA - European Environment Agency, 2011: Corine Land Cover. In <http://www.eea.europa.eu/publications/COROLandcover>, (Datum des letzten Besuches: 09 Juni 2011).
- Eidenbenz, C., Kaeser, C., Baltsavias, E., 2000: ATOMI – Automated reconstruction of topographic objects from aerial images using vectorized map information. In *IntArchPhRS*, vol. 33, no B3, S. 462–471.
- EN ISO 9000, 2005: Qualitätsmanagementsysteme- Grundlagen und Begriffe. In *ISO 9000:2005*.
- Fleiss, J. L., 1983: Statistical Methods for Rates and Proportions. In *John Wiley and Sons*, 2.Auflage, 223 Seiten.

- Foody, G.M., 2001: Status of land cover classification accuracy assessment. In *Remote Sensing of Environment*, vol. 80, S. 185–201.
- Foody, G.M., Mathur, A., 2004: A Relative Evaluation of Multiclass Image Classification by Support Vector Machines. In *IEEE Transactions on Geoscience and Remote Sensing*, vol. 42, no. 6, S. 1335–1343.
- Förstner, W., 1994: A framework for low level feature extraction. In *Computer Vision - ECCV, Proc. 5th ICCV*, vol. 2, S. 383–394.
- Fujimura, H., Ziems, M., Heipke, C., 2008: De-generalization of Japanese Road Data Using Satellite Imagery. In: *PFG*, vol. 5, S. 363–373.
- Fung, T., Chan, K.-C., 1994: Spatial Composition of Spectral Classes: A Structural Approach for Image Analysis of Heterogeneous Land-Use and Land-Cover Types. In *PE&RS*, vol. 60, no.2, S. 173–180.
- Gerke, M., Butenuth, M., Heipke, C., 2003: Automated update of road databases using aerial imagery and road construction data. In *IntArchPhRS*, vol 34, no. 3/W8, S. 99–104.
- Gerke, M., Heipke, C., 2008: Image based quality assessment of road databases. In *Int. J. of Geoinformation Science*, vol. 22, no. 8, S. 871–894.
- Gianinetto, M., 2008: Updating Large Scale Topographic Database in Italian Urban Areas with Submeter QuickBird Images. In *International Journal of Navigation and Observation*, Article ID 725429, 9 Seiten.
- Gimel'farb, G.L., 1996: Texture Modelling by Multiple Pairwise Pixel Interactions. In *IEEE TPAMI* 18, S. 1110–1114.
- Gimel'farb, G. L., 1997: Gibbs fields with multiple pairwise pixel interactions for texture simulation and segmentation. In *Rapport de recherche RR-3202, INRIA*, 70 Seiten.
- Gong, P., Mahler, S.A., Biging, G.S., Newburn, D. A., 2003: Vineyard identification in an oak woodland landscape with airborne digital camera imagery. In *Int. J. Remote Sensing*, vol 24, S. 1303–1315.
- Gonzalez, R., Woods, E., 2002: Digital Image Processing. In *Prentice Hall, Upper Saddle River (NJ)*, 2. Auflage, S. 622–624.
- Grenzdörffer, G., 2005: Flexible High Resolution Urban Remote Sensing with PFIFF – A Digital Low Cost System. In *IntArchPhRS*, vol 36, no 8/W27, auf CD.
- Grote, A., Heipke, C., 2008: Road extraction for the update of road databases in suburban areas. In *IntArchPhRS*, vol. 37, no. B3b, S. 563–568.
- Greve, W., Wentura, D., 1997: Wissenschaftliche Beobachtung: Eine Einführung. In *PVU/Beltz*, 185 Seiten.
- Hall, A., Louis, J., Lamb, D., 2003: Characterising and mapping vineyard canopy using high-spatial-resolution aerial multispectral images. In *Computer & Geosciences*, vol. 19, S. 813–822.
- Hall, A., Louis, J., Lamb, D., 2008: Low-resolution remotely sensed images of winegrape vineyards map spatial variability in planimetric canopy area instead of leaf area index. In *Australian journal of grape and wine research*, vol. 14, no. 1, S. 9–17.
- Hall-Beyer, M., 2008: The GLCM Tutorial. In <http://www.fp.ucalgary.ca/mhallbey/tutorial.htm> (Datum des letzten Besuches: 11. August 2011).
- Hanson, E., Wolff, E., 2010: Change Detection for Update of Topographic Databases through Multi-Level Region-Based Classification of VHR Optical and SAR Data. In *IntArchPhRS*, vol. 38, no 4-C7, auf CD.
- Haralick, R.M., Shanmugam, K., Dinstein, 1973: Texture features for image classification. In *IEEE Transactions on Systems, man. And Cybernetics*, SMC-3, S. 610–622.
- Heinert, M., 2010: Support Vector Machines – Teil 1: Ein theoretischer Überblick. In *zfv*, vol. 3, S. 179–189.
- Heinert, M., Riedel, B., 2010: Support Vector Maschine – Teil 2: Praktische Beispiele und Anwendungen. In *zfv*, vol. 5, S. 308– 313.
- Heipke, C., Mayer, H., Wiedemann, C., Jamet, O., 1997: Evaluation of Automatic Road Extraction. In *IntArchPhRS*, vol. 32, no. 4W2, S. 151–160.

- (Helmholz et al., 2007a) Helmholz, P., Gerke, M., Heipke, C., 2007: Automatic discrimination of farmland types using IKONOS imagery. In *IntArchPhRS*, vol. XXXVI, no. 3/W49A, S. 81–86.
- (Helmholz et al., 2007b) Helmholz, P., Becker, C., Heipke, C., Müller, S., Ostermann, J., Pahl, M., Ziem, M., 2007: Semi-automatic verification of geodata for quality management and updating of GIS. In *IntArchPhRS*, vol. 36, no. 4/W54, S. 9–13.
- Helmholz, P., Rottensteiner, F., Fraser, C., 2008: Enhancing the automatic verification of cropland in high-resolution satellite imagery. In *IntArchPhRS*, vol. 37, no B4, S. 385–390.
- (Helmholz et al., 2010a) Helmholz, P., Rottensteiner, F., Heipke, C., 2010: Automatic quality control of cropland and grassland GIS objects using IKONOS Satellite Imagery. In *IntArchPhRS*, vol. 38, no 7/B, S. 275–280.
- (Helmholz et al., 2010b) Helmholz, P., Becker, C., Breitkopf, U., Büschenfeld, T., Busch, A., Grünreich, D., Heipke, C., Müller, S., Ostermann, J., Pahl, M., Vogt, K., Ziem, M.; 2010: Semiautomatic Quality Control of Topographic Reference Datasets. In *IntArchPhRS*, vol. 38, no. 4, auf CD.
- Hermosille, T., Díaz-Manso, J.M., Ruiz, L.A., Recio, J.A., Fernández-Sarria, A., Ferradans-Nogueira, P., 2010: Parcel-Based image classification as a decision-making supporting tool for the Land Bank of Galicia (Spain). In *IntArchPhRS*, vol. 38 –W9, S. 40–45, auf CD.
- Hirose, Y., Mori, M., Akamatsu Y., Li, Y., 2004: Vegetation Cover Mapping Using Hybrid Analysis of Ikonos Data. In *IntArchPhRS*, vol. 35, no B7, S. 286–289.
- Hoberg, T., Rottensteiner, F., 2010: Classification of settlement areas in remote sensing imagery using Conditional Random Fields. In *IntArchPhRS*, vol. 38, no 7A, S. 53–58.
- Hoberg, T., Müller, S., 2011: Multitemporal Crop Type Classification Using Conditional Random Fields and RapidEye Data. In *IntArchPhRS*, vol. 28, no. 4/W19, 7 Seiten, auf CD.
- Hoffmann, A., Smith, G., Hesse, S., Lehmann, F., 2000: Die Klassifizierung hochaufgelösender Daten: Ein per-parcel-Ansatz mit Daten des digitalen Kampersystems HRSC-A. In *PFG*, vol 8, S. 1–17.
- Holland, D.A., Boyd, D.S., Marshall, P., 2006: Updating topographic mapping in Great Britain using imagery from high-resolution satellite sensors. In *ISPRS Journal of Photogrammetry and Remote Sensing*, vol. 60, S. 212–223.
- Hsu, C.-W., Chang, C.-C., Lin, C.-J., 2010: A Practical Guide to Support Vector Classification. In *Technical Report; Department of Computer Science, National Taiwan University: Taipei, Taiwan, 2003*, Verfügbar auf: <http://www.csie.ntu.edu.tw/~cjlin/papers/guide/guide.pdf> (Letzter Zugriff: 16. Dezember 2010).
- Huang, C., Davis, L. S., Townshend, J. R. G., 2002: An assessment of support vector machines for land cover classification. In *Int. J. Remote Sensing*, vol. 23, no. 4, S. 725–749.
- ISO 19113, 1999: Quality principles. In *ISO/TC 211 Projects*, Verfügbar auf: <http://www.isotc211.org>.
- Itzerott, S. and Kaden, K., 2006: Ein neuer Algorithmus zur Klassifizierung landwirtschaftlicher Fruchtarten auf Basis spektraler Normkurven. In *PFG*, vol. 6, S. 509–517.
- Itzerott, S. and Kaden, K., 2007: Gütebewertung für die Klassifizierung landwirtschaftlicher Fruchtarten aus spektralen Normkurven. In *PFG*, vol. 2, S. 109–120.
- Jähne, B., 2002: Digital Image Processing. In *Springer Verlag*, 585 Seiten.
- Janssen, L. L. F., van Amsterdam, J.D., 1991. An Object Based Approach to the Classification of Remotely Sensed Images. In *Proc. IGARSS*, S. 2191–2195.
- Joos, G., 2000: Zur Qualität von objektstrukturierten Geodaten. In *Schriftenreihe Studiengang Geodäsie und Geoinformation*, Universität d. Bundeswehr München, Heft 66, 141 Seiten.
- Klein, L., 1999: Sensor and Data Fusion, Concepts and Applications. In *SPIE Optical Engineering Press*, 2.Auflage, 226 Seiten.
- Knudsen, T., Olsen, B., 2003: Automated change detection for updates of digital map databases. In *PE&RS*, vol. 75, S. 1289–1296.
- Kreßel, U.H.G., 2002: Pairwise classification and support vector machines. In *Advances in Kernel Methods: Support Vector Machine Learning*, Herausgeber: B.Burges, C.J.C. Smola, The MIT Press, S. 255–268.

- Krickel, B., 2010: Informationserhebung zur Aktualisierung von ATKIS und Freizeitkataster in Nordrhein-Westfalen. In *ZfV*, vol. 4, S. 240–246.
- Krummel, J.R., Gardner, R.H., Sugihara, G., O'Neill, V., Coleman, P.R., 1987: Landscape patterns in a disturbed environment. In *OIKOS*, vol. 48, no. 3, S. 321–324.
- Kumar, S. and Hebert, M., 2006: Discriminative Random Fields. In *Int. J. Computer Vision*, vol. 68, no 2, S.179–201.
- Leignel, C., Caelen, O., Debeir, O., Hanson, E., Leloup, T., Simler, C., Beumir, C., Bontempi, G., Warz'ee, N., Wolff, E., 2010: Detecting Man-Made Structure Changes to Assist Geographic Data Producers in Planing Their Update Strategy. In *IntArchPhRS*, vol. 38, no 4-8-2-W9, auf CD.
- Leite, P.B.C., Feitosa, R.Q., Fromaggio, A.R., Costa, G.A.O.P., Pakzad, K., Sanche, I.D.A., 2008: Hidden Markov Models Applied in Agricultural Crops Classification. In *IntArchPhRS*, vol. 38, no 4, auf CD.
- Li, S., 2009: Markov Random Field Modeling in Image Analysis. In *Computer Science Workbench*, 3. Auflage, 362 Seiten.
- Lillesand, T.M., Kiefer, R.W., 2000: Remote Sensing and Image Interpretation. In *John Wiley & Sons Inc.*, 784 Seiten.
- Liu, Y., Zhen, Y.F., 2005: One-Against-All Multi-Class SVM Classification Using Reliability Measures. In *IEEE International Joint Conference on Neural Networks*, vol. 2, S. 849–854.
- Löcherbach, T., 1994: Fusion of multi-sensor images and digital map data for the reconstruction and interpretation of agricultural land-use units. In *IntArchPhRS*, vol. 15, no 3, S. 505–511.
- Lu, D. und Weng, Q., 2007: A survey of image classification methods and techniques for improving classification performance. In *Int. J. of Remote Sensing*, vol. 28, no. 5, S. 823–870.
- Lu, J., Plantaniotis, K. N., Venetsanopoulos, A. N., 2001: Face Recognition Using Feature Optimization and v-Support Vector Learning. In *Proc. IEEE International Workshop on Neural Networks for Signal Processing*, S. 373–382.
- Lucchese, L., Mitra, S.K., 2001: Volor image segmentation: a state-of-the-art survey. In *Image Processing, Vision and Pattern Recognition*, vol. 67, no. A(2), S. 207–221.
- Luo, J., Guo, C., 2003: Perceptual grouping of segmented regions in color images. In *Pattern Recognition*, vol. 36, S. 2781–2792.
- MacQueen, J.B., 1967: Some Methods for Classification and Analysis of MultiVariate Observations. In *Proc. of the 5th Berkeley Symposium on Mathematical Statistics and Probability*, vol. 1, S. 281–297.
- Marçal, A.R.S., Cunha, M., 2007: Vineyard monitoring in Portugal using multi-sensor satellite images. In *Proc. of the 27th EARSeL Symposium*, Netherlands, S. 327–335.
- Matikainen, I., Hyypä, J., Ahoka, E., Markelin, L., Kaarinen, H., 2010: Automatic Detection of Buildings and Changes in Buildings for Updating of Maps. In *Remote Sensing*, vol. 2, no. 5, S. 1217–1248.
- McLachlan, G. J., 2004: Discriminant Analysis and Statistical Pattern Recognition. In *Wiley-Interscience*, 552 Seiten.
- Müller, S., 2007. Extraktion baulich geprägter Flächen aus Fernerkundungsdaten zur Qualitätssicherung flächenhafter Geobasisdaten. In *ibidem-Verlag*, 152 Seiten.
- Müller, S., Heipke, C., Pakzad, K., 2010: Classification of farmland using multitemporal aerial images. In *IntArchPhRS*, vol. 38, no 4-8-2/W9, S. 70–74, auf CD.
- Nemmour, H., Chibani, Y., 2006: Multiple support vector machines for land cover change detection: An application for mapping urban extensions. In *ISPRS Journal of Photogrammetry & Remote Sensing*, vol. 61, S. 125–133.
- Olsen, B.P., 2004. Automatic Change Detection for Validation of Digital Map Databases. In *IntArchPhRS*, vol. 34, no B2, S. 569–574.
- Platt, J., 2000: Probabilistic outputs for support vector machines and comparison to regularized likelihood methods. In: *Advances in Large Margin Classifiers*, Herausgeber: A. Smola, P. Bartlett, B. Schölkopf and D. Schuurmans, Cambridge, MA, MIT Press, 11 Seiten.
- Peled, A., Gilichinsky, M., 2010: Knowledge-Based Classification of Land Cover for the Quality Assessment of GIS database. In *IntArchPhRS*, vol. 38, no 4-8-2-W9, S. 217–222.

- Perona, P., Malik, J., 1990: Scale-space and edge detection using anisotropic diffusion. In *IEEE TPAMI*, vol. 12, no. 7, S. 629–639.
- Poulain, V., Inglada, J., Spigai, M., 2008: High Resolution Remote Sensing Image Analysis with Exogenous Data: A Generic Framework. In *Proc. IGARSS*, vol. 2, S. 1025–1028.
- Poulain, V., Inglada, J., Spigai, M., Tournieret, J.-Y., Marthon, P., 2011: High-Resolution Optical and SAR Image Fusion for Building Database Updating. In *IEEE Transactions on Geoscience and Remote Sensing*, vol. 99, S. 1–11.
- Quinlan, J., R., 2003: C4.5 – Programs for Machine Learning. In *Morgan Kaufmann Publishers*, 302 Seiten.
- Rabatel, G., Delenne, C., Deshayes, M., 2008: A non-supervised approach using Gabor filters for vine-plot detection in aerial images. In *Computers and Electronics in Agriculture*, vol. 62, S. 159–168.
- Ranchin, T., Naert, B., Albuisson, M., Boyer, G., Astrand, P., 2001: An automatic method for vine detection in airborne imagery using wavelet transform and multiresolution analysis. In *PE&RS*, vol. 67, S. 91–98.
- Recio, J.A., Hermosilla, T., Ruiz, L.A., Fernández-Sarrià, A., 2011: Historical Land Use as a Feature for Image Classification. In *PE&RS*, vol. 77, no. 4, S. 377–387.
- Rengers, N., Prinz, T., 2009: JAVA-basierte Texturanalyse mittels Neighborhood Gray-Tone Differenz Matrix (NGTDM) zur Optimierung von Landnutzungsklassifikation in hoch auflösenden Fernerkundungsdaten. In *PFG*, vol. 5, S. 455–467.
- Römer, C., Plümer, L., 2010: Identifying Architectural Style in 3D City Models with Support Vector Machines. In *PFG*, vol. 5, S. 371–384.
- Rottensteiner, F., 2008: Automated updating of building data bases from digital surface models and multi-spectral images. In *IntArchPhRS*, vol. 37, no B3A, S. 265–270.
- Ruiz, L.A., Fernández-Sarrià, A., Recio, J.A., 2004: Texture feature extraction for classification of remote sensing data using wavelet decomposition: A comparative study. In *IntArchPhRS*, vol. 35, no B, S. 1109–1115.
- Ruiz, L.A., Recio, J.A., Hermosilla, T., 2007: Methods for automatic extraction of regularity patterns and its application to object-oriented image classification. In *IntArchPhRS*, vol. 36, no 3/W49A, S. 117–121.
- Russ, J.C., 1995: The Image Processing Handbook. In *CRC Press*, 2. Auflage, 674 Seite.
- Saarinen, K., 1994: Color Image Segmentation by a Watershed Algorithm and Region Adjacency Graph Processing. In *Proc. ICIP*, vol. 3, S. 1021–1025.
- Schölkopf, B., Sung, K., Burges, C., Girosi, F., Niyogi, P., Poggio, T., Vapnik, V., 1996: Comparing Support Vector Machine with Gaussian Kernels to Radial Basis Function Classifiers. In *A.I. Memo*, no 1599, MIT Press, 7.Seiten.
- Smith, G.M., Fuller, R.M., 2001: An integrated approach to land cover classification: an example in the Island of Jersey. In *Int. J. of Remote Sensing*, vol. 22, S. 3123–3142.
- Smits, C., Annoni, A., 1999: Updating Land-Cover Maps by Using Texture Information from Very High-Resolution Space-Borne Imagery. In *IEEE Transactions on Geoscience and Remote Sensing*, vol. 37, S. 1244–1254.
- Somers, B., Delalieux, S., Verstraeten, W. W., Coppin, P., 2009: A Conceptual Framework for the Simultaneous Extraction of Sub-pixel Spatial Extent and Spectral Characteristics of Crops. In *PE&RS*, vol. 75, no. 1, S. 57–68.
- Steinwart, I., Christmann, A., 2008: Support Vector Machines. In *Springer Verlag*, 601 Seiten.
- Statistisches Bundesamt, 2011: Statistisches Bundesamt Deutschland. In <http://www.destatis.de/jetspeed/portal/cms/> (Datum des letzten Besuches: 22. Juli 2011).
- Straub, B.-M., Wiedemann, C., Heipke, C., 2000: Towards the automatic interpretation of images for GIS update. In *IntArchPhRS*, vol. 33, S. 521–532.
- Sutton, R.N. und Hall, E.L., 1972: Texture measures for automatic classification of pulmonary disease. In *IEEE Transactions on Computers*, vol. 21, no. 7, S. 667–676.
- Tönnies, K. D., 2005: Grundlagen der Bildverarbeitung. In *Pearson–Studium-Verlag*, 341 Seiten.
- Trias-Sanz, R., 2006: Texture Orientation and Period Estimator between Forests, Orchards, Vineyard, and Tilled Fields. In *IEEE Transactions on Geoscience and Remote Sensing*, vol. 44, no. 10, S. 2755–2760.

- Unger, J., 2009: Untersuchung von Linien- und Kantenextraktionsalgorithmen im Rahmen der Verifikation von Ackerlandobjekten. *Bachelorarbeit am Institut für Photogrammetrie und GeoInformation* (http://www.ipi.uni-hannover.de/fileadmin/institut/pdf/Abschlussarbeiten/Bachelorarbeit_Jakob_Unger_16.12.09.pdf; Datum des letzten Besuches: 24. Februar 2012), 41 Seiten.
- Vapnik, V.N, 1998: Statistical Learning Theory. In *Wiley-Verlag*, 736 Seiten.
- Vosselmann, G., 1996: Uncertainty in GIS Supported Road Extraction. In *IntArchPhRS*, vol. 31, no B3, S. 906–916.
- Wackernagel, H., 2003: Multivariate Geostatistics. In *Springer Verlag*, 3. Auflage, 387 Seiten.
- Walter, V., 2004: Object-based classification of remote sensing data for change detection. In *ISPRS Journal of Photogrammetry and Remote Sensing*, vol. 58, S. 225–238.
- Warner, T.A. and Steinmaus, K., 2005: Spatial Classification of Orchards and Vineyards with High Spatial Resolution Panchromatic Imagery. In *PE&RS*, vol. 71, no.2, S. 179–187.
- Wassenaar, T., Robbez-Masson, J.-M., Andrieux, P., 2002: Vineyard identification and description of spatial crop structure by per field frequency analysis. In *Int. J. Remote Sensing*, vol. 23, no. 17, S. 3321–3325.
- Wilkinson, G. G., 1996: A review of current issues in the integration of GIS and remote sensing data. In *Int. J. Remote Sensing*, vol. 10, S. 85–101
- (Woodcock et al., 1988a) Woodcock, C.E., Strahler, A.H., Jupp, D.L.B., 1988: The Use of Variogram in Remote Sensing I: Scene Models and Simulated Images. In *Remote Sensing of Environment*, vol. 25, S. 323–348.
- (Woodcock et al., 1988b) Woodcock, C.E., Strahler, A.H., Jupp, D.L.B., 1988: The Use of Variogram in Remote Sensing II: Real Digital Images. In *Remote Sensing of Environment*, vol. 25, S. 349–379.
- Zhang, J., Goodchild, M., 2002: Uncertainty in Geographical Information. In *Taylor & Francis*, 266 Seiten.
- Ziems, M., Heipke, C., Rottensteiner, F., 2011: SVM-Based Road Verification with partly Non-Representative Training Data. In: *Proc. JURSE - Joint Urban Remote Sensing Event*, S. 37–40.

Abkürzungsverzeichnis

AdV	Arbeitsgemeinschaft der Vermessungsverwaltungen der Länder der Bundesrepublik Deutschland
ATKIS	Deutsches Amtlich Topographisch-Kartographisches Informationssystem
B	Blau (bzgl. des Kanales eines multispektralen Bildes)
BA	Bildanalyse
Basis-DLM	Digitale Basis-Landschaftsmodell
BKG	Bundesamt für Kartographie und Geodäsie
C5	C5-Klassifikationsverfahren
CLC	Europäisches CORINE Land Cover (GIS)
CRF	Conditional Random Fields
DGM	Digitalen Geländemodells
DGM _{Norm}	Normalisierten Digitalen Geländemodells
DLM-DE	Digitales Landschaftsmodell für Deutschland
DOM	Digitalen Oberflächenmodells
FN	False Negatives (Falsch negativ Entscheidung)
FP	False Positives (Falsch positiv Entscheidung)
G	Grün (bzgl. des Kanales eines multispektralen Bildes)
GIS	Geoinformationssystem
GLCM	Grey Level Co-Occurrence Matrix
GLDH	Gray Level Difference Histogram
IPI	Institut für Photogrammetrie und GeoInformation, Leibniz Universität Hannover
kNN	k-Nearest-Neighborhood
LDF	Linear Discriminant Function
MD	Minimum-Distance
MGCP	Internationales GIS des <i>Multinational Geospatial Co-production Program</i>
ML	Maximum-Likelihood
MRF	Markov Random Fields
NDVI	Normalized Differenced Vegetation Index
NGTDM	Neighborhood Gray-Tone Difference Matrix
NIR	Nahinfrarot (bzgl. des Kanales eines multispektralen Bildes)
R	Rot (bzgl. des Kanales eines multispektralen Bildes)
RAG	Region Adjacency Graph (Regionennachbarschaftsgraph)
SVM	Support Vector Machine
TG	Thematische Genauigkeit
TN	True Negatives (Richtig negativ Entscheidung)
TNT	Institut für Informationsverarbeitung, Leibniz Universität Hannover
TP	True Positives (Richtig positiv Entscheidung)
WiPKA-QS	Wissensbasierter Photogrammetrischer-Kartographischer Arbeitsplatz zur Qualitätssicherung

Lebenslauf / Curriculum vitae

Persönliches

Petra Helmholz
Geboren am 25. April 1979 in Halberstadt

Beruflicher Werdegang

seit 02/2006

Wissenschaftliche Mitarbeiterin,
Institut für Photogrammetrie und GeoInformation,
Leibniz Universität Hannover

05/2009 – 09/2009
11/2007 – 06/2008

Forschungsaufenthalte,
Cooperative Research Centre for Spatial Information (CRC SI),
University of Melbourne, Australien

11/2003 – 12/2005

Vermessungsreferendarin,
Landesamt für Vermessung und Geoinformation,
Land Sachsen-Anhalt

06/2001 – 10/2003

Tutorin,
Institut für Geodäsie und Geoinformationstechnik,
Technische Universität Berlin

1998 – 2000

Praktikantin bzw. freiberufliche Mitarbeiterin,
Amt für Landwirtschaft und Flurneuordnung Mitte in Halberstadt,
Land Sachsen-Anhalt

Ausbildung

11/2003 – 12/2005

Referendariat des Vermessungswesens, Land Sachsen-Anhalt
Abschluss: Assessor

10/1997 – 10/2003

Studium der Geodäsie, Technische Universität Berlin
Abschluss: Diplom – Ingenieur
Diplomarbeit (in Zusammenarbeit mit der Bauhaus Universität Weimar):
„Einsatz von Ultraschall bei hydrostatischen Messverfahren“

09/1993 – 09/1997

Gymnasium „Martineum“, Halberstadt
Abschluss: Abitur

09/1985 – 09/1993

Grundschule: POS „Ernst Thälmann“, Halberstadt

Dank

Dieses abschließende Kapitel möchte ich gerne nutzen, um einigen Menschen einen ganz herzlichen Dank auszusprechen.

Mein besonderer Dank gilt

meinem Doktorvater, Prof. Dr.-Ing. habil Christian Heipke, für die Möglichkeit der Anfertigung der vorliegenden Arbeit. Dazu gehören die fachliche Betreuung, die intensiven und konstruktiven Diskussionen und die Möglichkeiten des Austausches mit anderen Wissenschaftlern.

Prof. Dr.-Ing. Markus Gerke, für die Übernahme der Betreuung der Arbeit sowie für die Einführung in das Projekt WiPKA-QS und die Anleitung zum wissenschaftlichen Arbeiten in unserer gemeinsamen Zeit am IPI.

ganz speziell Privatdozent Dr. techn. Franz Rottensteiner, für die intensive Betreuung während der Ausarbeitung der Arbeit. Nur auf Grund der zahlreichen und immer sehr konstruktiven Diskussionen „flutscht“ das Dokument.

Prof. Clive Fraser, für die Übernahme des Gutachtens und die Betreuung meiner Arbeiten während meiner Zeiten in Melbourne.

dem BKG, für die Finanzierung des Projektes WiPKA-QS, in dessen Rahmen diese Arbeit entstand.

meinen Kollegen am IPI und TNT, für die positive und freundschaftliche Atmosphäre, die durch gegenseitige Unterstützung geprägt war.

meinen Freunden, die für die notwendige Ablenkung gesorgt sowie mich in schwierigen Zeiten immer unterstützt und aufgebaut haben.

meiner Familie, die mich stets in all meinen Bestrebungen und in jeglicher Hinsicht unterstützt hat und mich immer motivierte, meine Ziele zu erreichen.

and of course Phil, for his patience and all the sacrifices he made.