

Sabine Daniela Bauer

**Automatische Detektion
von Krankheiten auf Blättern von Nutzpflanzen**

München 2013

**Verlag der Bayerischen Akademie der Wissenschaften
in Kommission beim Verlag C. H. Beck**

ISSN 0065-5325

ISBN 978-3-7696-5130-0

**Diese Arbeit ist gleichzeitig veröffentlicht in:
Schriftenreihe des Instituts für Geodäsie und Geoinformation
der Rheinischen Friedrich-Wilhelms Universität Bonn
ISSN 1864-1113, Nr. 38, Bonn 2012**



Automatische Detektion von Krankheiten auf Blättern von Nutzpflanzen

Inaugural-Dissertation zur
Erlangung des akademischen Grades
Doktor-Ingenieur (Dr.-Ing.)
der Hohen Landwirtschaftlichen Fakultät
der Rheinischen Friedrich-Wilhelms Universität
zu Bonn

vorgelegt

am 15.07.2011 von

Sabine Daniela Bauer

aus Bonn

München 2013

Verlag der Bayerischen Akademie der Wissenschaften
in Kommission beim Verlag C. H. Beck

ISSN 0065-5325

ISBN 978-3-7696-5130-0

Diese Arbeit ist gleichzeitig veröffentlicht in:
Schriftenreihe des Instituts für Geodäsie und Geoinformation
der Rheinischen Friedrich-Wilhelms Universität Bonn
ISSN 1864-1113, Nr. 38, Bonn 2013

Adresse der Deutschen Geodätischen Kommission:



Deutsche Geodätische Kommission

Alfons-Goppel-Straße 11 • D – 80 539 München

Telefon +49 – 89 – 23 031 1113 • Telefax +49 – 89 – 23 031 -1283 / - 1100

e-mail hornik@dgfi.badw.de • <http://www.dgk.badw.de>

Diese Publikation ist als pdf-Dokument veröffentlicht im Internet unter der Adresse /
This volume is published in the internet

<<http://dgk.badw.de>> • <<http://hss.ulb.uni-bonn.de/2011/2744/2744-print.pdf>>

Prüfungskommission

Referent: Prof. Dr.-Ing. Dr. h.c. mult. Wolfgang Förstner

Korreferent: Prof. Dr. rer.nat. Lutz Plümer

Tag der mündlichen Prüfung: 02.12..2011

© 2013 Deutsche Geodätische Kommission, München

Alle Rechte vorbehalten. Ohne Genehmigung der Herausgeber ist es auch nicht gestattet,
die Veröffentlichung oder Teile daraus auf photomechanischem Wege (Photokopie, Mikrokopie) zu vervielfältigen

ISSN 0065-5325

ISBN 978-3-7696-5130-0

Danksagung

An dieser Stelle möchte ich mich ganz herzlich bei all meinen Freunden, Verwandten und Kollegen für die große Unterstützung bei meinem Promotionsvorhabens bedanken.

Insbesondere möchte ich Prof. Dr.-Ing. Dr.h.c.mult. Wolfgang Förstner für die zahlreichen kritischen und inspirierenden Diskussionen sowie die Ermöglichung der Teilnahme an Summer Schools, Konferenzen und anderen Weiterbildungsveranstaltungen ganz außerordentlich danken. Ich werde die Zeit in der Photogrammetrie in sehr angenehmer Erinnerung behalten.

Für die Betreuung der Arbeit aus phytomedizinischer Sicht sowie die vielen Tipps und Hinweise danke ich herzlich PD Dr. Erich-Christian Oerke und PD Dr. Ulrike Steiner.

Dr. Anne-Katrin Mahlein danke ich vielmals für die Anzucht, Inokulation und Pflege der Zuckerrübenpflanzen.

Heidi Hollander gilt mein großer Dank für die Korrektur meiner englischen Zeitschriftenartikel und Präsentationsfolien.

Die finanzielle Unterstützung des Promotionsvorhabens erfolgte durch das Graduiertenkolleg 722 "Einsatz von Informationstechniken zur Präzisierung des Pflanzenschutzes" der Deutschen Forschungsgemeinschaft (DFG), wofür ich mich ebenfalls herzlich bedanken möchte. Das agrarwissenschaftliche Lehrangebot des Graduiertenkollegs war für mich als Informatikerin zudem äußerst hilfreich bei der Einarbeitung in das Thema. Ebenfalls bedanken möchte ich mich bei den Mitgliedern des Graduiertenkollegs für die vielen fruchtbaren Diskussionen und gemeinsamen Exkursionen.

Für die Übernahme der Aufgabe des Korreferenten danke ich vielmals Prof. Dr. Lutz Plümer.

Der größte Dank gilt meinen Eltern. Denn ihre immerwährende Unterstützung in allen Lebenslagen und bestmögliche Förderung seit Kindesbeinen an hat diese Dissertation überhaupt erst möglich gemacht.

Zusammenfassung

Automatische Detektion von Krankheiten auf Blättern von Nutzpflanzen

In dieser Arbeit wird ein neuartiges, automatisches Verfahren zur Detektion von Krankheiten auf Blättern von Nutzpflanzen vorgestellt. Die Anwendung von Mustererkennungsverfahren zur Detektion von Blattkrankheiten ist noch ein sehr junges Forschungsgebiet. Ein einsetzbares Verfahren zur automatischen Detektion von Blattkrankheiten würde aber z.B. im Präzisionspflanzenbau zur Reduzierung der Spritzmittelmenge von großem Nutzen sein.

Das in dieser Arbeit vorgeschlagene Verfahren besteht aus einem hierarchischen Klassifikationsprozess. Im ersten Schritt dieses Prozesses wird eine pixelweise, adaptive Bayesklassifikation durchgeführt. Die adaptive Bayesklassifikation ist dabei speziell auf die Erkennung von Blattkrankheiten hin optimiert worden. Im zweiten Schritt erfolgt eine Glättung des pixelweisen Klassifikationsergebnisses durch eine Majoritätsfilterung oder mit Hilfe von bedingten Markoffschen Zufallsfeldern. Welche Vor- und Nachteile die beiden Glättungsarten bieten, wird in dieser Arbeit diskutiert. Der letzte Schritt des hierarchischen Klassifikationsprozesses besteht aus einer regionenbasierten Maximum-Likelihood-Klassifikation.

Eine besondere Herausforderung bei der Entwicklung des Klassifikationsverfahrens lag in der Berücksichtigung von Blattkrankheiten mit sehr kleinen Blattflecken. Um auch für diese Blattkrankheiten hohe Klassifikationsgenauigkeiten zu erzielen, wurde eine geeignete Bewertungsstrategie zur Beurteilung der Klassifikationsergebnisse entworfen und zur Optimierung des Klassifikationsprozesses verwendet.

Das entwickelte Klassifikationsverfahren basiert auf Bildern von einer RGB - und einer Multispektral (MS) - Kamera. Die Bilder der beiden Kameras wurden über ein 3D-Modell des jeweiligen Blattes fusioniert. Das zur Erstellung der 3D-Modelle entwickelte Verfahren bietet dabei neben der Sensorfusion auch das Potential zur Bestimmung der Lage und Orientierung von Blättern im Feld.

Im Rahmen der Entwicklung des Klassifikationsprozesses wurden weitere Fragestellungen bearbeitet. So wurde u.a. untersucht, inwieweit eine Erweiterung des Merkmalsvektors um Nachbarschaftsinformationen das Klassifikationsergebnis beeinflusst und ob Segmentierungsverfahren wie der Wasserscheidenalgorithmus zur Erkennung von Blattflecken eingesetzt werden können. Desweiteren wurde die Übertragbarkeit des Verfahrens auf andere Nutzpflanzen und Blattkrankheiten erforscht.

Abstract

Automatic detection of leaf diseases on agricultural crop

This thesis presents a novel, automatic method to detect leaf diseases on agricultural crop. The use of pattern recognition methods to detect leaf diseases is a very young research area. But an applicable method for the automatic detection of leaf diseases would be of great benefit for example in the area of precision agriculture.

The proposed method consists of an hierarchical classification process. In the first step of this process a pixelwise, adaptive Bayes classification is executed. Thereby the adaptive Bayes classifier is optimised for the detection of leaf diseases. In the second step a smoothing of the pixelwise result is carried out using a majority filter or a conditional random field (CRF). The advantages and disadvantages of the two ways of smoothing are the subject of a discussion in this thesis. The last step of the classification process consists of a regionbased Maximum-Likelihood classification.

A particular challenge in the development of the classification method lies in the consideration of leaf diseases with very small leaf spots. To achieve high classification accuracies for these leaf diseases, too, an adequate rating strategy was formulated and used for the optimisation of the classification process.

The developed classification method is based on images of a RGB - and a Multispectral (MS) - camera. The images of the two cameras was fused using a 3D-model of each single leaf. The prepared technique for the generation of the 3D-models of the leaves provides additionally the opportunity to determine the position and the orientation of the leaves in the field.

Within the framework of the development of the classification process further questions were examined. So it was investigated, in how far an enhancement of the feature vector with neighbourhood information affects the classification result and if it is possible to apply segmentation methods like the watershed algorithm for the recognition of leaf spots. In addition the transferability of the developed classification method to other plants and leaf diseases was proofed.

Inhaltsverzeichnis

1	Einleitung	13
1.1	Motivation und Zielsetzung	13
1.2	Bisherige Arbeiten	14
1.3	Herausforderungen und Arbeitshypothesen	16
1.4	Übersicht über die Arbeit	18
1.5	Notation	20
2	Grundlagen der verwendeten Klassifikationsverfahren	22
2.1	Ziel der Klassifikation	22
2.2	k -Nächste-Nachbarn Klassifikation	23
2.3	Bayesklassifikation	24
2.3.1	Bayesklassifikation mit Risikominimierung	25
2.3.2	Maximum <i>a-posteriori</i> Klassifikation	26
2.3.3	Maximum-Likelihood Klassifikation	26
2.4	Bedingte Markoffsche Zufallsfelder	27
2.5	Bewertung von Klassifikationsergebnissen	27
3	Verfahren zur Detektion von Blattkrankheiten	29
3.1	Versuchsaufbau	29
3.1.1	Pflanzenanzucht und Inokulation	29
3.1.2	Multimodales Kamerasystem	30
3.2	Fusion von Farb- und Multispektralbildern	31
3.2.1	Gemeinsame Orientierung von Farb- und Multispektralaufnahmen	31
3.2.2	Erstellung eines multispektralen Orthobildes	39
3.3	Hierarchischer Klassifikationsprozess	41
3.3.1	Übersicht über das Gesamtverfahren	42
3.3.2	Test- und Referenzdaten	43
3.3.3	Bewertungsstrategie bei der Klassifikation von Blattkrankheiten	44
3.3.4	Pixelweise, adaptive Bayesklassifikation	46
3.3.5	Glättung der pixelweisen Klassifikation zur Regionenbildung . .	47
3.3.6	Regionenbasierte Maximum-Likelihood Klassifikation	49

4	Experimente	51
4.1	Empirische Überprüfung der geometrischen Genauigkeit der Sensorfusion	51
4.2	Entwicklung der pixelweisen Klassifikationsstrategie	56
4.2.1	Verwendete Test- und Referenzdaten	56
4.2.2	Berücksichtigung der 4-er Nachbarschaft in den Merkmalen . . .	57
4.2.3	Optimierung der Bayesklassifikation	58
4.2.4	Bestimmung der Anzahl der Verteilungen pro Klasse	60
4.2.5	Bestimmung der <i>a-priori</i> Wahrscheinlichkeiten	63
4.2.6	Aufstellung der Gewichtungsfunktion für die Bayesklassifikation	65
4.3	Glättung der pixelweisen Klassifikation	71
4.3.1	Anwendung eines Majoritätsfilters oder Markoffscher Zufallsfelder	71
4.3.2	Vergleich der beiden Glättungsarten	72
4.4	Regionenbasierte Klassifikation	81
4.4.1	Einsatz von Wasserscheiden zur Regionenbildung	81
4.4.2	Maximum-Likelihood Klassifikation der geglätteten Regionen . .	84
4.4.3	Ergebnis der regionenbasierten ML-Klassifikation	85
4.5	Untersuchung der Übertragbarkeit des Verfahrens	89
4.5.1	Datenmaterial	89
4.5.2	Durchführung der Klassifikation	90
4.5.3	Ergebnis der Klassifikation	90
5	Schlussfolgerung	95

1 Einleitung

1.1 Motivation und Zielsetzung

Das Ziel dieser Arbeit ist die automatische Erkennung von Blattkrankheiten anhand von Multispektralbildern. Die Erkennung und Behandlung von Blattkrankheiten ist sowohl im Nutz- als auch im Zierpflanzenanbau von besonderer Wichtigkeit. Im Nutzpflanzenanbau führen Blattkrankheiten zu Ernteausschlag. Dies ist entweder darauf zurückzuführen, dass die Pflanzen ungenießbar werden, wie z.B. Salat, oder dass durch die zerstörte Blattfläche weniger Photosynthese betrieben werden kann und damit die Früchte der Pflanze nicht mit genügend Nährstoffen versorgt werden können. Dies ist z.B. bei der Zuckerrübe der Fall, wo bei einer Blattkrankheit weniger Zucker im Rübenkörper eingelagert wird oder sogar zum Austreiben neuer Blätter abgebaut wird. Im Zierpflanzenanbau führen Blattkrankheiten ebenfalls zu wirtschaftlichen Einbußen, da der Kunde eine makellose Pflanze erwartet. Durch Blattkrankheiten entstellte Pflanzen stören das ästhetische Empfinden und sind daher unverkäuflich. Aus diesen Gründen ist eine rechtzeitige Erkennung der spezifischen Blattkrankheit ganz wesentlich, damit eine entsprechende Behandlung der Pflanze eingeleitet werden kann.

Die korrekte Zuordnung der sichtbaren Symptome auf einem Blatt zu den jeweiligen auslösenden Pathogenen kann häufig nur von Experten, wie z.B. Phytomedizinern, erfolgen. Diese manuelle Zuordnung von Blattkrankheiten durch visuelle Inspektion der Blätter durch Experten ist aber teuer, langsam und hängt von der Erfahrung des jeweiligen Experten ab. Zudem kann das Ergebnis unterschiedlicher Experten durchaus verschieden ausfallen. Selbst das Ergebnis desselben Experten kann von einem Tag auf den anderen variieren (Bell *et al.*, 2002).

Eine automatische Erkennung der Blattkrankheiten durch Computeralgorithmen ist im Gegensatz dazu deutlich billiger, objektiv und für ein Blatt immer identisch. Es ist in diesem Fall kein Experte mehr notwendig, um die Blattkrankheit einer Pflanze zu bestimmen. Es wäre sogar denkbar, dass jedermann mit Hilfe einer Digicam Fotos einer betroffenen Pflanze macht und anschließend am Computer mit einer automatischen Erkennungssoftware für Blattkrankheiten das zugrundeliegende Pathogen bestimmt. Zudem kann durch die automatische Auswertung der Bilddaten die Menge der getesteten Pflanzen bei geringen Kosten deutlich erhöht werden.

Dies wäre z.B. im Bereich des Präzisionspflanzenschutzes von großem Vorteil. Denn das Ziel des Präzisionspflanzenschutzes ist es, nur die von einer Krankheit befallenen Areale eines Feldes zu behandeln und nicht, wie derzeit üblich, das komplette Feld

unabhängig von der Befallsverteilung (Rösch *et al.*, 2006). Die Voraussetzung dafür ist aber, Kenntnis über die Verteilung der Krankheit im Feld und evt. sogar der Verteilung innerhalb einer Pflanze bzw. auf einem Blatt zu besitzen. Die Gewinnung dieser Kenntnisse anhand einer manuellen Erfassung der Verteilung wäre extrem zeitintensiv, so dass die großflächige Anwendung des Präzisionspflanzenschutzes nur durch automatische Erfassungsverfahren möglich ist. Dazu könnten beispielsweise vorne am Traktor Kameras installiert werden. Während der Fahrt über das Feld würden dann die Pflanzen aufgenommen, die Bilder anhand der Erkennungssoftware für Blattkrankheiten automatisch analysiert und zuletzt gegebenenfalls entsprechend des Analyseergebnisses zielgerichtet behandelt werden.

Die Entwicklung von für dieses Szenario notwendigen automatischen Erkennungsverfahren von Blattkrankheiten ist Ziel dieser Arbeit.

1.2 Bisherige Arbeiten

Publikationen zur Thematik der automatischen Erkennung von Blattkrankheiten basierend auf Bildinformationen gibt es bisher kaum. Darunter sind die folgenden vier Arbeiten von Cui *et al.* (2010), Sanyal und Patel (2008), Huang (2007) und Pydipati *et al.* (2006). Bei diesen Arbeiten basiert die Klassifikation auf der Farbinformation der Kanäle Rot, Grün und Blau. Tatsächlich lässt sich aber die gesunde, chlorophyllhaltige und die durch biotische oder abiotische Umwelteinflüsse zerstörte, nicht-chlorophyllhaltige Blattfläche am ehesten im nahen Infrarotbereich voneinander separieren. Dies liegt an der hohen Reflexivität des Chlorophylls in diesem Wellenlängenbereich und der Tatsache, dass die zerstörte Blattfläche in der Regel kein Chlorophyll mehr enthält. In der Fernerkundung wird dieses Wissen schon seit langem zur Unterscheidung zwischen Vegetation und Nicht-Vegetation eingesetzt (Albertz, 1999). Ebenso werden im Fernerkundungsbereich zur Einschätzung der Vitalität von Pflanzen Multispektral- oder sogar Hyperspektralbilder eingesetzt, welche auch den Infrarotbereich mit abdecken (siehe z.B. Lelong *et al.* (1998) oder Mewes *et al.* (2010)). Aus diesem Grunde berücksichtigt das in dieser Dissertation entwickelte Verfahren zur automatischen Detektion von Blattkrankheiten nicht nur die Farbinformationen der Kanäle Rot, Grün und Blau, sondern zusätzlich auch die Infrarotinformation.

Die verwendeten Klassifikationsmethoden und untersuchten Pflanzen und Krankheiten differieren in allen oben genannten Arbeiten. So setzen Pydipati *et al.* (2006) eine Diskriminanzanalyse zur Identifikation von drei Blattkrankheiten bei Grapefruit-Gehölzen ein. Sowohl Huang (2007) als auch Sanyal und Patel (2008) verwenden beide künstliche neuronale Netze. Huang (2007) setzen das künstliche neuronale Netz aber zur Erkennung von Blattkrankheiten von Orchideen-Setzlingen der Gattung *Phalaenopsis* ein, während Sanyal und Patel (2008) in ihrer Arbeit eine Mustererkennungsmethode basierend auf einem mehrlagigen Perzeptrons zur Detektion von zwei Krankheiten auf Reispflanzen einführen. Die neueste Arbeit von Cui *et al.* (2010) untersucht die Er-

kennung einer Rostinfektion auf Sojabohnenblätter. Diese Arbeit von Cui *et al.* (2010) stellt eine manuelle Schwellwertmethode und einen Ansatz für eine automatische Rosterkennung mit Hilfe einer Analyse des Schwerpunkts der Verteilung der Farbinformationen im polaren Koordinatensystem vor.

Die vorliegende Dissertation prüft nun eine Kombination aus pixelweiser Bayesklassifikation, Glättung mit bedingten Markoffschen Zufallsfeldern und anschließender regionenbasierter Maximum-Likelihood Klassifikation auf ihre Eignung zwei verschiedene Krankheiten auf Zuckerrübenblättern zu detektieren. Bei der Entwicklung dieses Verfahrens wird für jedes Blatt ein 3D-Modell erstellt. Die 3D-Struktur ist erforderlich, da die verwendeten Kameras nicht in allen vier Kanälen empfindlich sind. Die aus diesem Grunde notwendige Fusionierung der Sensordaten von der RGB- und Multispektralkamera erfolgt über ein 3D-Modell des Blattes. Das in dieser Arbeit angewendete Verfahren zur Erstellung der 3D-Modelle bietet zudem potentiell auch die Möglichkeit, es zukünftig zur Lagebestimmung der Blätter im Feld einzusetzen. Bei allen vier obigen Publikationen spielte dieser Aspekt der Gewinnung einer 3D-Struktur von Blätter bzw. Pflanzen dagegen keine Rolle.

Unabhängig von der Klassifikation von Blattkrankheiten existieren bereits Publikationen, die sich mit der Gewinnung der 3D-Struktur von Pflanzen beschäftigen. In diesen Arbeiten wurden verschiedene Technologien zur Erstellung des 3D-Modells eingesetzt. Stuppy *et al.* (2003) analysierten die dreidimensionale Struktur von Pflanzen mit Hilfe von hochaufgelösten Röntgenaufnahmen. Dieses Verfahren ist durch den Einsatz von Röntgenstrahlen aber nicht im Feld einsetzbar. Hanan *et al.* (2004) setzten ein Laserscanning-Verfahren ein. Die Erstellung von Laserscans ist allerdings teuer und zudem problematisch bei Wind. Denn während der Aufnahme von Objekten mittels eines Laserscanners sollten sich diese nicht bewegen, da die aufgenommene 3D-Information ansonsten inkonsistent ist. Ein deutlich billigeres Verfahren zur Erstellung von 3D-Modellen basiert auf Stereobildern. Zur Vermeidung von Inkonsistenzen durch die Bewegung der Blätter besteht hier die Möglichkeit die Stereobilder synchron aufzunehmen. Pan *et al.* (2004) erforschten als erstes die Eignung dieses bildbasierten Verfahrens zur 3D-Rekonstruktion von Pflanzen. Sie führten allerdings noch eine manuelle Selektion der Zuordnungen durch. Das in dieser Arbeit für die Erstellung von 3D-Modellen von Blättern angepasste Verfahren funktioniert dagegen vollautomatisch.

Im Computergraphik-Bereich wird deutlich intensiver an der Erstellung von realistischen 3D-Modellen von Pflanzen geforscht. Hier dient die Rekonstruktion dazu natürliche Modelle von Pflanzen für virtuelle Umgebungen zu erhalten. Diese 3D-Modelle spiegeln aber nicht unbedingt exakt die 3D-Struktur der jeweiligen Pflanze wieder (Ma und Zha, 2009). Eine Übersicht über den Stand der Technik der bildbasierten Modellierung von Pflanzen und Bäumen im Computer Vision Bereich bietet das Buch von Kang und Quan (2010).

1.3 Herausforderungen und Arbeitshypothesen

Die automatische Detektion von Blattkrankheiten bietet einige Herausforderungen. So ist z.B. wegen der großen Zahl an Nutzpflanzen, welche in Form und Farbe der Blätter variieren können, und der Vielzahl u. U. gleichzeitig auftretender Blattkrankheiten, eine Klassifizierung der einzelnen Blattkrankheiten häufig nicht ganz einfach. Um diese Vielfalt etwas einzugrenzen, basiert die Entwicklung des Verfahrens zur automatischen Detektion von Blattkrankheiten zunächst nur auf einer Nutzpflanze. Anschließend wird dann die Übertragbarkeit des Verfahrens auf andere Nutzpflanzen getestet. Als Nutzpflanze zur Verfahrensentwicklung dient die Zuckerrübe, da sie eine hohe wirtschaftliche Bedeutung in Deutschland und USA besitzt (Wolf und Verreet (2002), Steddom *et al.* (2005)). Die Übertragbarkeit des Verfahrens wird anhand von Reisblättern geprüft.

Die nächste Herausforderung bei der Erkennung von Blattkrankheiten liegt in der Veränderung der visuellen Symptome einer Blattkrankheit im Laufe ihrer Krankheitsentwicklung. Die Blätter der Zuckerrübe können beispielsweise u.a. von dem Braunrost (*Uromyces betae*), dem Mehltau (*Erysiphe betae*) und den Blattfleckenkrankheiten (*Cercospora beticola* sowie *Ramularia beticola*) befallen werden. Nach KWS (2010) sind bei der Blattfleckenkrankheit *Cercospora beticola* die Flecken z.B. anfangs annähernd rund und gleichmäßig dunkelgrau bis schwarz. Beim Fortschreiten der Krankheit wird das Kreisinere dann hell, während sich außen ein dunkelbrauner Rand bildet. Die Größe der Flecken beträgt dabei ungefähr 2-3 mm. Im weiteren Verlauf der Krankheit fließen die einzelnen Flecken häufig zusammen und schließlich stirbt das komplette Blatt ab.

In Abbildung 1.1 sind Flecken verschiedener Entwicklungsstadien dieser Blattfleckenkrankheit visualisiert. Zwei Flecken im Anfangsstadium sind durch hellblaue Ellipsen markiert.



Abbildung 1.1: Ausprägung der Blattfleckenkrankheit *Cercospora beticola* auf einem Zuckerrübenblatt. Die hellblau markierten Flecken befinden sich noch im Anfangsstadium.

Ein weiteres Beispiel zur Veränderung des Krankheitsbildes bietet der Braunrost *Uromyces betae*. Er entwickelt sich nach Hillnhütter und Mahlein (2008) zunächst nur innerhalb des Blattes. Typische Symptome von *Uromyces betae* sind im Anfangsstadium daher lediglich kleine chlorotische Stellen. Erst im fortgeschrittenen Stadium bricht das Blatt stellenweise auf, um die braunen Rostsporen abzugeben. Diese sind dann als kleine ungefähr 1 mm große dunkelbraune Punkte zu sehen (siehe Abbildung 1.2).

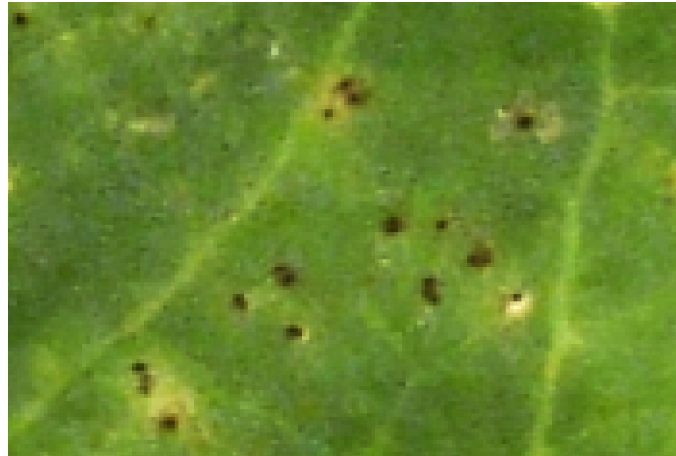


Abbildung 1.2: Ausprägung des Braunrostes *Uromyces betae* auf einem Zuckerrübenblatt.

Die unterschiedlichen Blattflecken auf einer bestimmten Nutzpflanze können also sowohl auf die verschiedenen Blattrankheiten dieser Nutzpflanze zurückgeführt werden als auch auf die verschiedenen Krankheitsstadien einer bestimmten Blattrkrankheit. Um die Bandbreite dieser verschiedensten Blattflecken auf den zu klassifizierenden Zuckerrübenblättern etwas einzuschränken, erfolgte die Entwicklung des Verfahrens anhand der Blattfleckenkrankheit *Cercospora beticola* und des Braunrostes *Uromyces betae*. Die Blattfleckenkrankheit *Cercospora beticola* gilt als die bedeutendste pilzliche Erkrankung der Zuckerrübe in Mitteleuropa (Wolf und Verreet, 1997). Die wirtschaftliche Bedeutung des Braunrostes ist zwar eher gering (KWS, 2010), dafür bietet er aber einige Herausforderung bei der Klassifikation. Dies ist auf die geringe Größe der Rostpusteln zurückzuführen. Durch ihre Kleinheit werden die Rostpusteln z.B. bei der Anwendung von Glättungsfiltren eliminiert oder bei Einsatz von Segmentierungsverfahren im Zuge der Rauschunterdrückung entfernt. Zudem ist bei einem nicht so ausgeprägten Befall der Zuckerrübe mit dieser Blattrkrankheit die *a-priori* Wahrscheinlichkeit unter Umständen nicht höher als 1/2 Promille. Dies führt dann bei Klassifikatoren, welche die Gesamtklassifikationsgenauigkeit optimieren, ebenfalls zu Problemen (siehe Kapitel 3.3.3).

Weitere Herausforderungen liegen bei einer Anwendung im Feld im Auftreten von Verdeckungen, Spiegelungen und der Sichtbarkeit von Blattober- bzw. Blattunterseiten im Bild. Die beiden Blattseiten weisen dabei in der Regel eine unterschiedliche

Struktur auf. Ebenso sind die visuellen Symptome einer Blattkrankheit häufig von der jeweiligen Blattseite abhängig. In dieser Arbeit wurde diese Problematik vermieden, indem die Aufnahmen der Blätter im Labor erfolgte. Dort wurden die Blätter dann jeweils einzeln fotografiert (siehe Kapitel 3.1.2). Beeinträchtigungen der Farbinformation durch Reflexionen oder Schattenwurf auf der eher glänzenden und krausen Oberfläche der Zuckerrübenblättern wurden ebenfalls durch eine diffuse Beleuchtung der Blätter während der Aufnahme verhindert.

Basierend auf den RGB- und Multispektralbildern der Zuckerrübenblättern mit Symptomen der beiden Blattkrankheiten *Cercospora beticola* und *Uromyces betae* in verschiedenen Entwicklungsstadien wurden in dieser Arbeit folgende globalen Arbeitshypothesen verifiziert:

1. Die Berücksichtigung der 4-er Nachbarschaft in den Merkmalen erhöht die Klassifikationsgenauigkeiten signifikant.
2. Der Bayesklassifikator kann durch die Anpassung der Likelihood- und Gewichtungsfunktion sowie der *a-priori* Wahrscheinlichkeiten für die Klassifikation von seltenen Klassen, wie dem Braunrost *Uromyces betae*, optimiert werden.
3. Die Glättung eines pixelweisen Klassifikationsergebnisses mit bedingten Markoffschen Zufallsfeldern ist einer Glättung durch einen Majoritätsfilter überlegen.
4. Eine regionenbasierte Klassifikation der segmentierten Objekte erhöht die Klassifikationsgenauigkeiten aller Klassen.

1.4 Übersicht über die Arbeit

Nach der bis jetzt erfolgten Motivation des Themas und der Auseinandersetzung mit den auf diesem Gebiet bisher erschienen Publikationen sowie den Herausforderungen bei der Lösung dieser Aufgabe, folgt nun nach der in dieser Arbeit verwendeten Notation zuerst das Theorie-Kapitel 2. Dieses Kapitel beschreibt die Grundlagen der in dieser Arbeit verwendeten Klassifikationsverfahren. Neben der Erläuterung des Ziels einer Klassifikation und der Bewertung von Klassifikationsergebnisses stellt dieses Kapitel den k -Nächsten-Nachbarn Klassifikator, die Bayesklassifikation und die bedingten Markoffschen Zufallsfelder vor.

Auf diesen Grundlagen aufbauend beschreibt das nächste Kapitel 3 das in dieser Arbeit entwickelte Verfahren zur Detektion von Blattkrankheiten basierend auf Farb- und Multispektralbildern. Dazu erfolgt zunächst eine Beschreibung des Versuchsaufbaus anhand dessen die Bilder der Pflanzen aufgenommen wurden. Anschließend wird die Fusion der Farb- und Multispektralbilder mit Hilfe einer gemeinsamen Orientierung der Aufnahmen und der Erstellung eines multispektralen Orthobildes erläutert. Das letzte Unterkapitel geht schließlich auf den in dieser Dissertation erarbeiteten hierarchischen

Klassifikationsprozess ein. Dieser Klassifikationsprozess besteht aus einer pixelweisen, adaptiven Bayesklassifikation, einer anschließenden Glättung des pixelweisen Ergebnisses mit Hilfe eines Majoritätsfilter oder mit bedingten Markoffschen Zufallsfeldern und einer im letzten Schritt durchgeführten regionenbasierten Maximum-Likelihood Klassifikation.

Die zur Entwicklung dieses Verfahrens durchgeführten Experimente stellt das nächste Kapitel 4 vor. Das erste Experiment dient zur empirischen Überprüfung der geometrischen Genauigkeit der Sensorfusion. Das nächste Unterkapitel enthält die zur Entwicklung der pixelweisen Klassifikationsstrategie notwendigen Experimente. Anschließend folgen die Experimente zur Glättung der pixelweisen Klassifikation und zur regionenbasierten Klassifikation. Als letztes wird das Experiment zur Übertragbarkeit des Verfahrens auf durch Eisentoxizität geschädigte Reisblätter vorgestellt.

Das letzte Kapitel 5 erläutert schließlich die Schlussfolgerungen, welche aus der Arbeit gezogen werden können.

1.5 Notation

Symbol Bedeutung

Allgemein

x, X	Skalar
$\boldsymbol{x}, \boldsymbol{X}$	Vektor
X	Matrix
$\mathbf{X}$	homogene Matrix
$\mathcal{X}$	Menge
$\hat{x}$	geschätzter Wert
$\tilde{x}$	wahrer Wert

Reservierte Symbole

$p(\dots)$	Wahrscheinlichkeitsdichte
$P(\dots)$	Wahrscheinlichkeit
R	Rotationsmatrix
K	Kamera-Kalibriermatrix
P	Kamera-Projektionsmatrix
$\boldsymbol{t}$	Translationsvektor
$\mathcal{S}$	Menge aller Elemente, wie z.B. Pixel
s	Index eines Elements
S	Anzahl der Elemente
$\boldsymbol{y}$	Alle Beobachtungen des Bildes, es gilt $\boldsymbol{y} = [y_0, y_1, \dots, y_s, \dots, y_{S-1}]^T$
y_s	Beobachtung an dem Element s
$\boldsymbol{h}_s(\boldsymbol{y})$	Merkmale an dem Element s basierend auf den Beobachtungen $\boldsymbol{y}$

$\mathbf{x}$	Komplette Ground Truth Klassenkonfiguration, es gilt $\mathbf{x} = [x_0, x_1, \dots, x_s, \dots, x_{S-1}]^T$
x_s	Ground Truth Klasse des Elements s
$\hat{\mathbf{x}}$	Komplette geschätzte Klassenkonfiguration
$\hat{x}_s$	Klassifikationsergebnis des Elements s
L	Anzahl der Klassen
$\mathcal{L}$	Menge der Klassen
l	Index einer Klasse

 Statistische Symbole

σ	Standardabweichung
Σ	Kovarianzmatrix

 Mathematische Operatoren

$\mathbf{X}^T$	Transponierte
$\mathbf{X}^{-1}$	Inverse einer Matrix

2 Grundlagen der verwendeten Klassifikationsverfahren

2.1 Ziel der Klassifikation

Das Ziel einer Klassifikation von Bildern ist es, jedes Element s , wie z.B. ein Pixel oder eine Region, aus der Menge $\mathcal{S} = \{0, \dots, s, \dots, S - 1\}$ eines Bildes genau einer Klasse $\hat{x}_s$ zuzuweisen (Li, 2001). Im vorliegenden Fall sind beispielsweise folgende Klassen $\mathcal{L} = \{0, \dots, l, \dots, L - 1\}$ vertreten:

- 0: gesunder Blattbereich
- 1: mit *Cercospora beticola* infizierter Blattbereich
- 2: mit *Uromyces betae* infizierter Blattbereich

Die Klassifikation eines Elementes s erfolgt anhand seines Merkmalvektors $\mathbf{h}_s(\mathbf{y})$. Der Merkmalsvektor basiert auf den Beobachtungen $\mathbf{y}$ des Bildes. Im einfachsten Falle besteht er nur aus den Beobachtungen y_s eines Elementes s selbst, d.h. es gilt $\mathbf{h}_s(\mathbf{y}) = y_s$. Beobachtungen y_s eines Elementes s können zum Beispiel folgende Farbinformationen sein:

- Rot
- Grün
- Blau
- Infrarot

Der Merkmalsvektor $\mathbf{h}_s(\mathbf{y})$ kann aber auch Nachbarschaftsinformationen mitberücksichtigen, wie zum Beispiel die Farbinformation des oberen, unteren, linken und rechten Nachbarn:

$$\mathbf{h}_s(\mathbf{y}) = \begin{bmatrix} y_s \\ y_l \\ y_r \\ y_o \\ y_u \end{bmatrix} \quad \text{mit } l = \text{links}, r = \text{rechts}, o = \text{oben}, u = \text{unten}$$

Zum Training eines Klassifikators und zur Kontrolle des Klassifikationsergebnisses werden Referenzdaten benötigt, welche als korrekt angenommen werden. Die Referenzdaten für das gesamte Bild werden im Folgenden mit $\mathbf{x}$ bezeichnet. Die Referenzklasse an einem Element s ist x_s , wobei $x_s \in \mathcal{L}$ mit $\mathcal{L} = \{0, \dots, l, \dots, L - 1\}$ gilt. Die vom Klassifikator geschätzte Klassenzuordnung eines Elementes s wird als $\hat{x}_s$ angegeben und entsprechend ist die komplette, geschätzte Klassenkonfiguration eines Bildes definiert als $\hat{\mathbf{x}}$.

Nachstehend werden verschiedene Klassifikationsverfahren vorgestellt. Jedes dieser Klassifikationsverfahren muss zunächst mit Hilfe von Trainingsdaten trainiert werden. Erst nach Abschluß dieses Lernprozesses kann der Klassifikator dann selbstständig neue Daten klassifizieren.

2.2 *k*-Nächste-Nachbarn Klassifikation

Der *k*-Nächste-Nachbarn (*k*NN) Klassifikator ordnet ein Merkmalsvektor $\mathbf{h}_s(\mathbf{y})$ derjenigen Klasse zu, zu der die Mehrheit der *k* nächsten Nachbarn gehört. Die Nähe wird dabei mit einer geeigneten Metrik gemessen, wie z.B. mit dem Euklidischen Abstand (Fukunaga, 1972). Training dieses Klassifikators bedeutet, dass eine Menge von Beobachtungen $\mathbf{y}$ und die dazugehörigen Referenzdaten $\mathbf{x}$ bekannt sind. Damit ist bei dem *k*NN Klassifikator keinerlei Wissen über die Verteilung der Daten notwendig.

Ein weiterer Vorteil dieses Klassifikators ist die Möglichkeit, den so genannten Bayesfehler ε^* zu schätzen. Der Bayesfehler ist ein Maß für die Separierbarkeit von Klassen. Denn bei bekannter Verteilung der Merkmale von zwei Klassen korrespondiert er zu der überlappenden Fläche der beiden Verteilungen. Mit dem Bayesfehler kann man daher bestimmen, welche Klassifikationsgenauigkeiten bei bestimmten Daten und Merkmalen erreicht werden können. Als Ergebnis bekommt man eine Spanne, in welcher die Klassifikationsergebnisse anderer Klassifikatoren liegen werden. Für diese Schätzung muss allerdings die Anzahl der Trainingsdaten groß genug sein, damit der zu klassifizierende Merkmalsvektor $\mathbf{h}_s(\mathbf{y})$ und seine nächsten Nachbarn möglichst nah im Merkmalsraum benachbart sind (Fukunaga, 1972).

Wenn die Fehlerwahrscheinlichkeit einer Klasse als ε und die Anzahl der Klassen als L definiert ist, berechnet sich der Bayesfehler ε^* nach Fukunaga (1972) folgendermaßen

$$\frac{L-1}{L} - \sqrt{\left(\frac{L-1}{L}\right)^2 - \varepsilon \frac{L-1}{L}} \leq \varepsilon^* \leq \varepsilon. \quad (2.1)$$

2.3 Bayesklassifikation

Eine Bayesklassifikation basiert, wie der Name schon andeutet, auf dem Bayes-Theorem

$$P(l \mid \mathbf{h}_s(\mathbf{y})) = \frac{p(\mathbf{h}_s(\mathbf{y}) \mid l) \cdot P(l)}{P(\mathbf{h}_s(\mathbf{y}))}. \quad (2.2)$$

In dieser Formel ist nach Duda *et al.* (2001)

- $p(\mathbf{h}_s(\mathbf{y}) \mid l)$ die bedingte Wahrscheinlichkeitsdichte für das Auftreten von $\mathbf{h}_s(\mathbf{y})$, falls das Element zur Klasse l gehört. Aus ihr ergibt sich die Likelihoodfunktion $L(\mathbf{h}_s(\mathbf{y}) \mid l)$.
- $P(l)$ die *a-priori* Wahrscheinlichkeit. Sie gibt an, mit welcher Wahrscheinlichkeit die Klasse l auftritt.
- $P(\mathbf{h}_s(\mathbf{y}))$ die Wahrscheinlichkeit für das Auftreten des Merkmalvektors $\mathbf{h}_s(\mathbf{y})$. Da dieser Term lediglich eine normalisierende Wirkung hat, kann er während der Klassifikation unberücksichtigt bleiben.
- $P(l \mid \mathbf{h}_s(\mathbf{y}))$ die *a-posteriori* Wahrscheinlichkeit. Sie gibt die Wahrscheinlichkeit für die Klasse l bei der Beobachtung des Merkmals $\mathbf{h}_s(\mathbf{y})$ an.

Die Bestimmung der *a-priori* Wahrscheinlichkeit $P(l)$ erfolgt z.B. anhand der Auftrittswahrscheinlichkeit der verschiedenen Klassen. In diesem Falle ergibt sich die *a-priori* Wahrscheinlichkeit bei einer Anzahl von S Elementen, von denen S_l zu der Klasse l gehören nach $P(l) = \frac{S_l}{S}$.

Die Likelihood-Funktion $L(\mathbf{h}_s(\mathbf{y}) \mid l)$ kann z.B. durch eine Mischverteilung aus Gauß-Verteilungen angenähert werden. Um die Parameter einer solchen Mischverteilung zu berechnen, bietet sich z.B. der Expectation-Maximization (EM) Algorithmus von Bilmes (1998) an. Für die Berechnung der Parameter werden dem Algorithmus N Trainingsdaten der jeweiligen Klasse l , welche auf den Beobachtungen $\mathbf{y}$ beruhen, und die Anzahl C an Clustern der Mischverteilung übergeben. Zusätzlich benötigt der Algorithmus für jedes Cluster c die vorab z.B. mittels des k -means Algorithmus' geschätzten Mittelwerte μ_c und *a-priori* Wahrscheinlichkeiten α_c sowie die Kovarianzmatrix Σ_c .

Im ersten Schritt, dem sogenannten E-Schritt, berechnet der Algorithmus die *a-posteriori* Wahrscheinlichkeit $P(c \mid \mathbf{y})$ für die Zugehörigkeit der Beobachtungen $\mathbf{y}$ zum Cluster c . Die Likelihood-Funktion $\Phi(\mathbf{y} \mid \mu_c, \Sigma_c)$ gibt die Wahrscheinlichkeitsdichte für das Auftreten der Beobachtungen $\mathbf{y}$ bei gegebenem Mittelwert μ_c und Varianz Σ_c an. Die Wahrscheinlichkeit $P(c \mid \mathbf{y})$ berechnet sich somit mit Hilfe der folgenden Formel

$$P(c | \mathbf{y}) = \frac{\alpha_c \Phi(\mathbf{y} | \mu_c, \Sigma_c)}{\sum_{c=1}^C \alpha_c \Phi(\mathbf{y} | \mu_c, \Sigma_c)}. \quad (2.3)$$

Im zweiten Schritt, dem sogenannten M-Schritt, wird die Mittelwertsmatrix μ_c , die *a-priori* Wahrscheinlichkeiten α_c und die Kovarianzmatrix Σ_c neu berechnet. Der E- und M-Schritt des EM-Algorithmus' wird solange wiederholt, bis die Änderungen klein genug bleiben.

2.3.1 Bayesklassifikation mit Risikominimierung

Eine Bayesklassifikation mit Risikominimierung berücksichtigt neben der Likelihood-Funktion und der *a-priori* Wahrscheinlichkeiten der einzelnen Klassen auch Kosten, welche bei den verschiedenen Entscheidungen entstehen. Für jede Entscheidungsmöglichkeit muss vorab die Höhe der entstehenden Kosten festgelegt werden. Diese werden in einer sogenannten Kostenmatrix eingetragen (siehe Tabelle 2.1).

<i>Kostenmatrix</i>					
%	$\hat{0}$	...	$\hat{l}$	...	$\widehat{L-1}$
$\tilde{0}$	$k(0, 0)$	...	$k(0, l)$	...	$k(0, L-1)$
	...	...	...	...	...
$\tilde{l}$	$k(l, 0)$	...	$k(l, l)$	...	$k(l, L-1)$
	...	...	...	...	...
$\widetilde{L-1}$	$k(L-1, 0)$	...	$k(L-1, l)$	...	$k(L-1, L-1)$

Tabelle 2.1: Kosten bei unterschiedlichen Entscheidungen.

Das Ziel einer Bayesklassifikation mit Risikominimierung ist die Minimierung dieser Kosten. D.h. der Klassifikator entscheidet sich für diejenige Klasse, die die größten Kosten verursachen würde, wenn man nicht sie sondern eine andere Klasse wählen würde.

Binärfall Nach Fukunaga (1972) fällt die Entscheidung für Klasse 0, wenn

$$\frac{k(1, 0) - k(1, 1)}{k(0, 1) - k(0, 0)} P(1) L(\mathbf{h}_s(\mathbf{y}) | 1) < P(0) L(\mathbf{h}_s(\mathbf{y}) | 0) \quad (2.4)$$

gilt, andernfalls wird sich für Klasse 1 entschieden. D.h. wir entscheiden uns für die Klasse 0, wenn die Kosten bei einer Fehlentscheidung bei Klasse 1 niedriger sind bzw. im Umkehrschluß bei Klasse 0 höher sind.

Mehrklassenfall Für den Mehrklassenfall wird die obige Formel folgendermaßen erweitert

$$\hat{x}_s = \arg \min_{x_s} \sum_{l=0}^L (k(l, x_s) - k(l, l)) P(l) L(\mathbf{h}_s(\mathbf{y}) \mid l). \quad (2.5)$$

2.3.2 Maximum *a-posteriori* Klassifikation

Die Maximum *a-posteriori* (MAP) Klassifikation ist ein Spezialfall der Bayesklassifikation mit Risikominierung. In diesem Fall wird davon ausgegangen, dass alle Fehlentscheidungen die gleichen Kosten verursachen. Daher kann hier die Kostenmatrix unbeachtet bleiben und man maximiert stattdessen lediglich das Produkt aus der Likelihood-Funktion und der *a-priori* Wahrscheinlichkeit

$$\hat{x}_s = \arg \max_{x_s} P(l) L(\mathbf{h}_s(\mathbf{y}) \mid l). \quad (2.6)$$

Die Voraussetzung zur Berechnung der *a-posteriori* Wahrscheinlichkeit ist, dass sowohl die Likelihood-Funktion als auch die *a-priori* Wahrscheinlichkeit bekannt sind. Diese müssen daher beim Training des Klassifikators mit Hilfe der Trainingsdaten bestimmt werden (Fink, 2003).

Nach Beendigung der Trainingsphase des Klassifikators und der damit geschätzten Likelihood-Funktion sowie der *a-priori* Wahrscheinlichkeiten $P(l)$ der einzelnen Klassen kann die Testphase beginnen. Für jeden Merkmalsvektor $\mathbf{h}_s(\mathbf{y})$ des Testdatensatzes wird die *a-posteriori* Wahrscheinlichkeit nach folgender Formel maximiert

$$\hat{x}_s = \arg \max_{x_s} P(l) \sum_{c=1}^C \alpha_c p(\mathbf{h}_s(\mathbf{y}) \mid \mu_c, \Sigma_c), \quad (2.7)$$

wobei C die Anzahl der Verteilungen pro Klasse, α_c die *a-priori* Wahrscheinlichkeit, μ_c den Mittelwert und Σ_c die Kovarianzmatrix der jeweiligen Verteilung c angibt.

2.3.3 Maximum-Likelihood Klassifikation

Die Maximum-Likelihood (ML) Klassifikation ist ein Sonderfall der MAP-Klassifikation. Hier wird für alle Klassen die gleiche *a-priori* Wahrscheinlichkeit angenommen. Daher wird bei der ML-Klassifikation, wie der Name schon sagt, nur die Likelihood-Funktion maximiert

$$\hat{x}_s = \arg \max_{x_s} L(\mathbf{h}_s(\mathbf{y}) \mid l). \quad (2.8)$$

In der Testphase wird die Klassenzugehörigkeit mit Hilfe folgender Formel berechnet

$$\hat{x}_s = \arg \max_{x_s} \sum_{c=1}^C \alpha_c p(\mathbf{h}_s(\mathbf{y}) \mid \mu_c, \Sigma_c). \quad (2.9)$$

2.4 Bedingte Markoffsche Zufallsfelder

Die bisher vorgestellten Klassifikationsverfahren klassifizieren jedes Element, wie z.B. ein Pixel oder eine Region, völlig unabhängig von der geschätzten Klassenzugehörigkeit der Nachbapixel oder -regionen. D.h. im Falle einer pixelweisen Klassifikation ist die komplette geschätzte Klassenkonfiguration $\hat{\mathbf{x}}$ definiert als $\hat{\mathbf{x}} = \{\hat{x}_s\}_{s \in \mathcal{S}}$.

Im Gegensatz dazu wird bei den hier vorgestellten Markoffschen Zufallsfeldern die komplette Klassenkonfiguration eines Bildes in einer einzigen globalen Entscheidung $\hat{\mathbf{x}}$ geschätzt, indem die optimale Entscheidung für alle Elemente $\hat{x}_s$ gemeinsam getroffen wird. Die bedingten Markoffschen Zufallsfelder wurden ursprünglich von Lafferty *et al.* (2001) zur Segmentierung und Labelling von 1-D Textsequenzen vorgeschlagen. Kumar und Hebert (2006) erweiterten das Konzept so, dass es auch auf 2-D Objekten, wie z.B. Bildern angewandt werden kann.

Das in dieser Dissertation angewandte Verfahren basiert auf den Ausarbeitungen in Bauer *et al.* (2011). Das Lernen der unbekannten Modellparameter der bedingten Markoffschen Zufallsfelder erfolgt darin im Prinzip nach einem Standard Maximum Likelihood Ansatz. Aufgrund der Komplexität der dabei zu lösenden Gleichung wird die Likelihood-Funktion allerdings durch eine Pseudolikelihood-Funktion basierend auf der Arbeit von Kumar *et al.* (2005) und Korč und Förstner (2008) angenähert.

Die anschließende Schätzung der Klassenkonfiguration erfolgt in Bauer *et al.* (2011) mittels Linearer Programmierung (LP) Relaxation basierend auf der Arbeit von Schlesinger (1976). Unabhängig von Schlesinger (1976) haben aber auch andere Arbeitsgruppen Lösungen mittels der LP-Relaxation vorgeschlagen (siehe Chekuri *et al.* (2001), Koster *et al.* (1998), Wainwright *et al.* (2005)).

Die Lösung der Linearen Programmierung wird in dieser Dissertation mittels des so genannten Active-Set Algorithmus berechnet. Der verwendete Algorithmus ist eine Abwandlung des Verfahrens von Gill *et al.* (1984) und Gill *et al.* (1991) (siehe The MathWorks (2010)).

2.5 Bewertung von Klassifikationsergebnissen

In den vorigen Kapiteln wurden verschiedene Verfahren zur Klassifikation von Bildern vorgestellt. Wenn nun das Klassifikationsergebnis vorliegt, soll die Qualität dieses Ergebnisses mit Hilfe von Referenzdaten beurteilt werden. Im optimalen Fall stimmen das Klassifikationsergebnis und die Referenzklassen überein, in der Regel liegen aber Abweichungen vor. Eine Übersicht über die Übereinstimmungen und Abweichungen bietet eine sogenannte Konfusionsmatrix (siehe Tabelle 2.2). In dieser Tabelle sind auf der linken Seite die Referenzklassen $\tilde{l}$ aufgeführt und oben die Testklassen $\hat{l}$. Die Wahrscheinlichkeit mit der ein Pixel aus der Klasse $\tilde{l}$ auch der Klasse $\hat{l}$ zugewiesen wird, ist definiert als $P(\hat{l} | \tilde{l})$. Optional ist es möglich auch eine Rückweisungsklasse r einzuführen. In diese Klasse kommen diejenigen Pixel, bei denen keine eindeutige Zu-

2 Grundlagen der Klassifikationsverfahren

ordnung des Elementes zu einer Klasse möglich ist, weil z.B. die Wahrscheinlichkeiten für zwei Klassen gleich groß sind.

Konfusionsmatrix						
%	$\hat{0}$	...	$\hat{l}$	...	$\widehat{L-1}$	r
$\tilde{0}$	$P(\hat{0} \tilde{0})$	...	$P(\hat{l} \tilde{0})$	...	$P(\widehat{L-1} \tilde{0})$	$P(r \tilde{0})$
$\tilde{l}$	...	...	...	...	...	...
	$P(\hat{0} \tilde{l})$	...	$P(\hat{l} \tilde{l})$	...	$P(\widehat{L-1} \tilde{l})$	$P(r \tilde{l})$
$\widetilde{L-1}$	...	...	...	...	...	...
	$P(\hat{0} \widetilde{L-1})$	...	$P(\hat{l} \widetilde{L-1})$	...	$P(\widehat{L-1} \widetilde{L-1})$	$P(r \widetilde{L-1})$

Tabelle 2.2: Aufbau einer Konfusionsmatrix.

Wenn $\tilde{S}_l$ die Anzahl der Pixel ist, welche korrekter Weise zur Klasse $\tilde{l}$ gehören und $\hat{S}_l$ die Anzahl der Pixel von $\tilde{S}_l$ ist, die der Klasse $\hat{l}$ bei der Klassifikation zugeordnet wurden, dann berechnet sich die Wahrscheinlichkeit $P(\hat{l} | \tilde{l})$ anhand von

$$P(\hat{l} | \tilde{l}) = \frac{\hat{S}_l}{\tilde{S}_l}. \quad (2.10)$$

Die Wahrscheinlichkeiten in einer Zeile der Konfusionsmatrix summieren sich wie folgt

$$\sum_{\hat{l}=0}^{\widehat{L-1}} P(\hat{l} | \tilde{l}) = 1. \quad (2.11)$$

Die Werte auf der Diagonalen von links oben nach rechts unten in der Konfusionsmatrix sind die Wahrscheinlichkeiten $P(\hat{l} | \tilde{l})$. Diese geben an, dass ein Pixel bei der Klassifikation seiner korrekten Klasse zugeordnet wurde. Diese Wahrscheinlichkeiten werden als Klassifikationsgenauigkeiten der jeweiligen Klasse $\tilde{l}$ bezeichnet. Die Bewertung eines Klassifikators erfolgt anhand der Höhe der erzielten Klassifikationsgenauigkeiten, welche möglichst maximal sein sollten (siehe Formel (2.12)). Entsprechend weist ein guter Klassifikator nur minimale Fehlerwahrscheinlichkeiten auf (siehe Formel (2.13)) (Bishop, 2007).

$$\max_l P(\hat{l} | \tilde{l}) \quad (2.12)$$

$$\min_{\hat{l}, \tilde{l}} \sum_{\tilde{l}} \sum_{\hat{l} \neq \tilde{l}} P(\hat{l} | \tilde{l}) \quad (2.13)$$

3 Verfahren zur Detektion von Blattkrankheiten

3.1 Versuchsaufbau

Das Verfahren zur Detektion von Blattkrankheiten ist anhand von Aufnahmen von Zuckerrübenblättern entwickelt worden, welche entweder mit der Blattfleckenkrankheit *Cercospora beticola* oder dem Braunrost *Uromyces betae* inokuliert wurden. Wie die Anzucht und Inokulation der Zuckerrübenpflanzen erfolgte, ist im ersten Unterkapitel beschrieben. Das zweite Unterkapitel erläutert das Aufnahmesystem, mit welchem die Bilder von den Blättern angefertigt wurden.

3.1.1 Pflanzenanzucht und Inokulation

Die Pflanzenanzucht und Inokulation der Zuckerrüben wurde durch den Bereich Phyto-medizin des Instituts für Nutzpflanzenwissenschaften und Ressourcenschutz (INRES) durchgeführt. Für die folgenden Experimente wurden 30 Zuckerrübenpflanzen der Sorte „Pauletta“ von der KWS GmbH in Einbeck in einem kommerziell verfügbaren Konzentrat von der Klasmann-Deilmann GmbH in Plastiktöpfen (Ø 13 cm) angezogen (siehe Mahlein *et al.*, 2010). Die Inokulation der gesunden Pflanzen erfolgte, sobald 4 Blätter voll entwickelt waren. Dies entspricht dem Wachstumsstadium 14 auf der BBCH Skala (Meier *et al.*, 1993). Es wurden 15 Pflanzen mit der Blattfleckenkrankheit *Cercospora beticola* und die anderen 15 mit dem Braunrost *Uromyces betae* inokuliert.

Nach Hillnhütter und Mahlein (2008) tritt die Blattfleckenkrankheit *Cercospora beticola* vorwiegend in Gebieten mit Temperaturen von 22–25°C und einer hohen relativen Luftfeuchtigkeit, Taubildung bzw. regelmäßigem Niederschlag auf. Als optimale Bedingungen für den Rost gelten nach Hillnhütter und Mahlein (2008) Temperaturen bis 20°C und eine relative Luftfeuchtigkeit von etwa 100%. Aus diesem Grund wurden die *Cercospora*-Pflanzen nach der Inokulation für 48 Stunden in eine Umgebung mit einer Luftfeuchtigkeit von 100% RH und einer Temperatur von 25/20 Grad Celsius gebracht, während die *Uromyces*-Pflanzen bei gleicher Luftfeuchtigkeit nur einer Temperatur von 19/16 Grad Celsius ausgesetzt wurden. Zur weiteren Inkubation wurden alle Pflanzen anschließend in eine Umgebung mit 23/20 Grad Celsius bei 60 +- 10% RH transferiert (Mahlein *et al.*, 2010).

Nähere Details zu den Anzuchtbedingungen der Pflanzen, zur Herstellung der Inoku-

lumslösung sowie dem Inokulationsvorgang können der Veröffentlichung von Mahlein *et al.* (2010) entnommen werden.

3.1.2 Multimodales Kamerasystem

Von jeder inokulierten Pflanze wurden alle zwei Tage jeweils zwei markierte Blätter fotografiert. Die Bildaufnahme erfolgte im Labor unter kontrollierter, diffuser Beleuchtung mit zwei verschiedenen Kameras. Die eine Kamera war die handelsübliche RGB-Kamera "FujiFilm FinePix S5600" mit einer Auflösung von 2592×1944 Pixeln. Bei der anderen Kamera handelte es sich um die Multispektral- (MS-) Kamera "Tetracam ADC" mit einer Auflösung von 1280×1024 Pixeln. Die MS-Kamera hat die Kanäle Rot, Grün und Infrarot (700 - 950 nm).

Wie schon in der Einleitung im Kapitel 1 erläutert, sollen im Folgenden als Klassifikationsmerkmale sowohl die Farbinformationen der RGB-Kamera als auch die Infrarotinformation der MS-Kamera verwendet werden. Die dazu notwendige Fusion der Farbinformationen der beiden Kameras erfolgte über ein 3D-Modell des jeweiligen Blattes. Die Erstellung der 3D-Modelle geschah wiederum über Stereo-Aufnahmen der Blätter. Aus diesem Grunde wurde jedes Blatt von vier unterschiedlichen Positionen aus mit der RGB-Kamera fotografiert. Der durchschnittliche Abstand der Aufnahmepositionen betrug ungefähr 10 cm bei einem Abstand von 30 cm zum Objekt. Zur Erfassung der Infrarotinformation erfolgte zusätzlich von jedem Blatt eine fünfte Aufnahme mit der MS-Kamera. Aus dem Aufnahmeabstand von 30cm resultiert im RGB-Bild eine Auflösung des Blattes von ungefähr 0.0967 mm und im MS-Bild von ungefähr 0.2123 mm.

In Abbildung 3.1 ist eine schematische Darstellung des Versuchsaufbaus zu sehen. Das zu fotografierende Blatt wurde auf einem kleinen Tischchen mit einem Gummiband am Blattansatz fixiert und dann von oben aufgenommen. In der Zeichnung ist die RGB-Kamera schwarz dargestellt und die MS-Kamera rot.

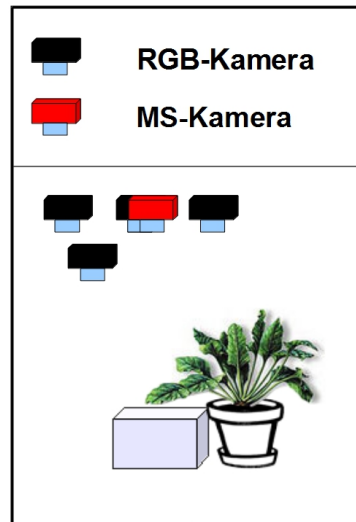


Abbildung 3.1: Schema des Versuchsaufbaus (Bauer *et al.*, 2009).

3.2 Fusion von Farb- und Multispektralbildern

Um bei der Klassifikation sowohl die Farbinformationen der RGB- als auch der MS-Kamera verwenden zu können, müssen die Farb- und Multispektralbilder fusioniert werden (siehe Abbildung 3.2). D.h. die Bilder müssen so übereinander projiziert werden, dass für jeden Punkt $\mathbf{P}_{xyz}$ auf dem Blatt sowohl die Farbinformation des RGB-Bildes als auch die Farbinformation des MS-Bildes bekannt ist.

Die Fusionierung ist hier anhand eines 3D-Modells des Blattes vorgenommen worden. Das 3D-Modell wiederum ist mit Hilfe der Stereo-Bilder erstellt worden. Zur Erstellung eines 3D-Modells müssen zunächst alle Aufnahmepositionen bestimmt werden. Dies wird auch als gemeinsame Orientierung bezeichnet. Wie diese gemeinsame Orientierung von Farb- und Multispektralaufnahmen durchgeführt werden kann, erläutert das folgende Unterkapitel. Nach erfolgreicher Erstellung des 3D-Modells kann ein sogenanntes multispektrale Orthobild erzeugt werden. Als Orthofoto wird ein rektifiziertes Bild bezeichnet, aus welchem geometrische Verzerrungen, bedingt z.B. durch die Krümmung eines Zuckerrübenblattes, eliminiert wurden. Das zur Erstellung der Orthofotos angewendete Vorgehen erläutert das zweite Unterkapitel.

3.2.1 Gemeinsame Orientierung von Farb- und Multispektralaufnahmen

Das Ziel der gemeinsamen Orientierung der Farb- und Multispektralaufnahmen ist es, die Lage der Kameras während der verschiedenen Aufnahmen sowie die Lage der Objektpunkte, d.h. des Blattes, im Raum zu bestimmen. Die Berechnung der Orientierung kann man in zwei Schritte untergliedern, zunächst die Berechnung der inneren Orientierung und anschließend die Berechnung der äußeren Orientierung.

Die innere Orientierung einer Kamera hängt von dem Aufbau der Kamera und des

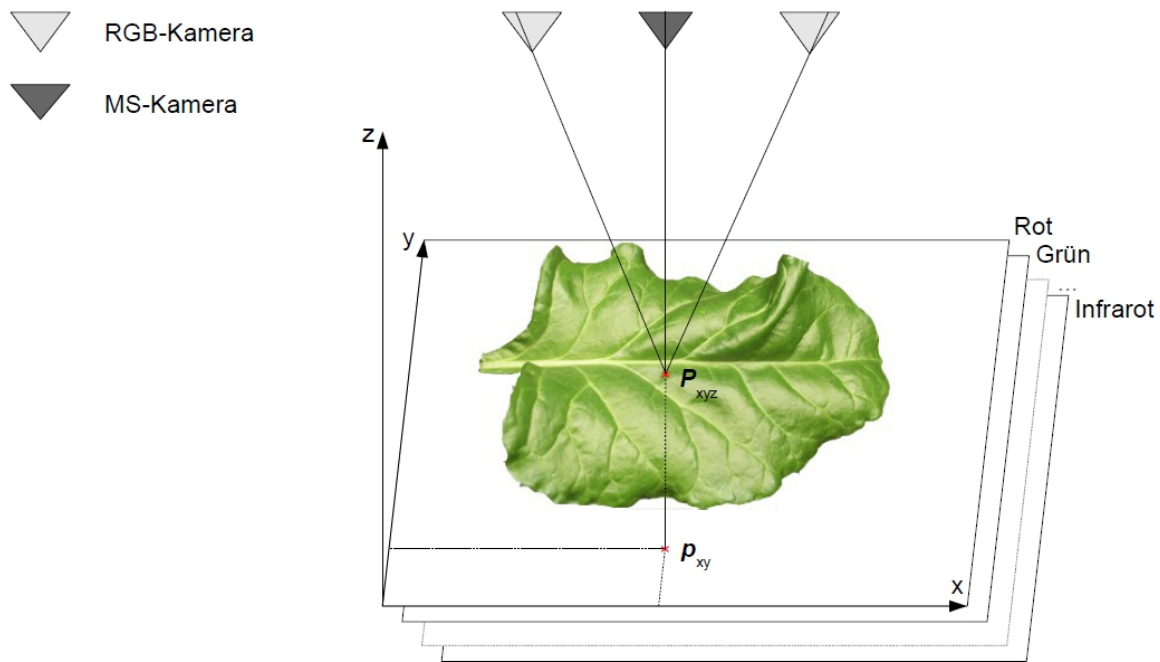


Abbildung 3.2: Fusion der Bilder von der RGB- und der MS-Kamera: Die Zeichnung zeigt das 3D-Modell eines Blattes mit den 3D-Oberflächenpunkten P_{xyz} . Für alle Punkte P_{xyz} ist die Farbinformation von jedem einzelnen Bild bekannt. In der Zeichnung sind zur Vereinfachung nur zwei RGB-Kameras eingezeichnet, in der Realität handelte es sich aber um vier Kameras. Dies bedeutet, dass für jeden 3D-Punkt P_{xyz} und damit auch für jeden 2D-Punkt p_{xy} 15 Farbwerte bekannt sind. Die 3D-Punkte P_{xyz} werden auf die unten dargestellte Referenzebene abgebildet, so dass ein rektifiziertes Bild mit den 2D-Punkten p_{xy} entsteht (siehe Bauer *et al.*, 2011).

Objektivs ab und ist daher von Kamera zu Kamera unterschiedlich. Sie muss in der Regel nur einmal bestimmt werden, solange man die Einstellungen der Kamera danach nicht mehr verändert. In der äußeren Orientierung wird die Lage der Kameras im Raum berechnet. Diese kann entweder absolut oder relativ bestimmt werden. Bei einer absoluten Orientierung wird ein Bezug zu einem realen Koordinatensystem hergestellt. Bei einer relativen Orientierung dagegen wird angenommen, dass die Kamera sich bei der ersten Aufnahme im Nullpunkt des Koordinatensystems befindet und es werden dann die Lage der anderen Aufnahmestandorte relativ zu diesem ersten Aufnahmestandort berechnet. Im Gegensatz zur inneren Orientierung muss die äußere Orientierung für jede Szene erneut berechnet werden. Für detailliertere Informationen siehe z.B. McGlone *et al.* (2004), Kraus (2004) und Luhmann (2003).

Innere Orientierung Bei der inneren Orientierung wird nach Hartley und Zisserman (2003) u.a. die Verschiebung des Hauptpunktes, die radiale Verzerrung durch die Linse und die Kamerakonstante berechnet. Als Hauptpunkt $\mathbf{H}$ wird der Fußpunkt des Lots vom Projektionszentrum $\mathbf{P}$ auf die Bildebene bezeichnet. Im Idealfall sollte sich der Hauptpunkt $\mathbf{H}$ im Zentrum der Bildebene befinden, in der Regel ist er aber bedingt durch den Kameraaufbau etwas versetzt. Das Projektionszentrum $\mathbf{P}$ liegt bei modernen Kameras im Linsensystem. Die Kamerakonstante c ist die Distanz von dem Projektionszentrum $\mathbf{P}$ zur Bildebene. Die radiale oder auch tangentielle Verzerrung des Bildes wird durch das Objektiv verursacht. Eine Durchführung der inneren Orientierung führt zur Korrektur der Lage des Hauptpunktes und der Entfernung von Verzerrungen aus dem Bild, so dass als Ergebnis ein ideales Bild zur Verfügung steht.

Die innere Orientierung wurde sowohl für die RGB- als auch für die MS-Kamera einmalig mit den Aufnahmeeinstellungen berechnet. Die Bestimmung der inneren Orientierung wird auch als Kamerakalibrierung bezeichnet. Zur Kamerakalibrierung werden mit der jeweiligen Kamera aus verschiedenen Aufnahmepositionen und Aufnahmegerichtungen sowie unterschiedlichen Drehwinkeln der Kamera um die Aufnahmeachse 24 Bilder von einem Testfeld erstellt. Auf dem Testfeld befinden sich auf zwei verschiedenen Ebenen Passmarken, deren Koordinaten bekannt sind. Somit können die 3D-Objektkoordinaten eindeutig den entsprechenden 2D-Bildkoordinaten zugeordnet werden und anhand dessen die Kameraparameter ermittelt werden. Die zur Kalibrierung eingesetzte Software wurde von Abraham (1999) entwickelt.

Äußere Orientierung Bei der äußeren Orientierung wird nach McGlone *et al.* (2004) die Position der Kamera im Raum bestimmt. D.h. es wird die Translation $\mathbf{t}$ und Rotation R der Kamera in Bezug auf ein gegebenes Weltkoordinatensystem berechnet, so dass schließlich die Lage des Projektionszentrums, die Aufnahmerichtung der Kamera und deren Drehwinkel um die Aufnahmeachse zum Aufnahmezeitpunkt vorliegt. Durch die äußere Orientierung der Kamera können die Objektkoordinaten im Weltkoordinatensystem in das Kamerakoordinatensystem umgerechnet werden. Diese Transformation

vom Weltkoordinatensystem in das Kamerakoordinatensystem hat sechs Freiheitsgrade, drei für die Rotation R und drei für die Translation t . Die Koordinaten des Punktes P_w im Weltkoordinatensystem können mit folgender Formel in die Koordinaten P_k im Kamerakoordinatensystem umgerechnet werden

$$P_k = R \cdot P_w + t.$$

Die Berechnung der Rotation und Translation für die äußere Orientierung erfolgte im vorliegenden Fall mit der Software AURELO (Läbe und Förstner, 2006). Die Software liefert als Ergebnis für jedes Bild eine sogenannte Projektionsmatrix. Eine Projektionsmatrix enthält die Rotations- und Translationsparameter der äußeren Orientierung sowie die Parameter der Kalibriermatrix der inneren Orientierung der jeweiligen Kamera. Die Software nimmt dabei an, dass der Ursprung des Weltkoordinatensystems im Projektionszentrum der ersten Aufnahme liegt. Desweiteren geht das Programm AURELO zur Zeit noch davon aus, dass alle Bilder mit derselben Kamera aufgenommen worden sind. Bei der Verwendung von zwei verschiedenen Kameras, wie es in dieser Arbeit mit dem Einsatz einer RGB- und MS-Kamera durchgeführt wurde, müssen die Bilder der einen Kamera daher so umgerechnet werden, als wenn sie mit der anderen Kamera aufgenommen worden wären. Am einfachsten geschieht dies dadurch, dass man für beide Kameras die innere Orientierung berechnet und damit die Verzerrungen und die Verschiebung des Hauptpunktes herausrechnet. Bei den so transformierten Bildern muss dann nur noch die Kamerakonstante aneinander angepasst werden. Da sowohl die Auflösung als auch die Kamerakonstante der RGB-Kamera größer ist als die der MS-Kamera, wurden die Bilder der MS-Kamera an die Bilder der RGB-Kamera angepasst. Wie dies konkret geschehen ist, erläutert der folgende Paragraph.

Anpassung der Kamerakonstante In Abbildung 3.3 ist eine schematische Zeichnung der Transformation der Kamerakonstante des MS-Bildes zu sehen. Sowohl das MS- als auch das RGB-Bild sind bereits entzerrt und die Lage der Hauptpunkte wurden korrigiert. Es führt somit eine Symmetrieachse vom oberen Objektpunkt durch das Zentrum des MS- und des RGB-Bildes. Die Kamerakonstante der MS-Kamera beträgt $c_i = 1365.71$ und die Kamerakonstante der RGB-Kamera $c_r = 2841.02$.

Passt man nun das MS-Bild an die Kamerakonstante der RGB-Kamera an, wird das Bild auf die RGB-Bildebene projiziert und vergrößert sich dabei auf 2663×2130 Pixel. Die Vergrößerung wurde mit Hilfe einer bilinearen Interpolation durchgeführt. Das RGB-Bild selbst hat aber nur eine Auflösung von 2592×1944 Pixeln. In der Zeichnung ist diese Fläche hellgrau dargestellt.

Für die Weiterverarbeitung mit AURELO müssen die Bilder allerdings die gleiche Auflösung aufweisen. Daher mussten entweder die Ränder vom MS-Bild beschnitten oder die Größe des RGB-Bildes angepasst werden. Um einen Datenverlust bei dem MS-Bild zu verhindern, ist das RGB-Bild vergrößert worden, indem ein schwarzer Rand um das Bild gelegt wurde.

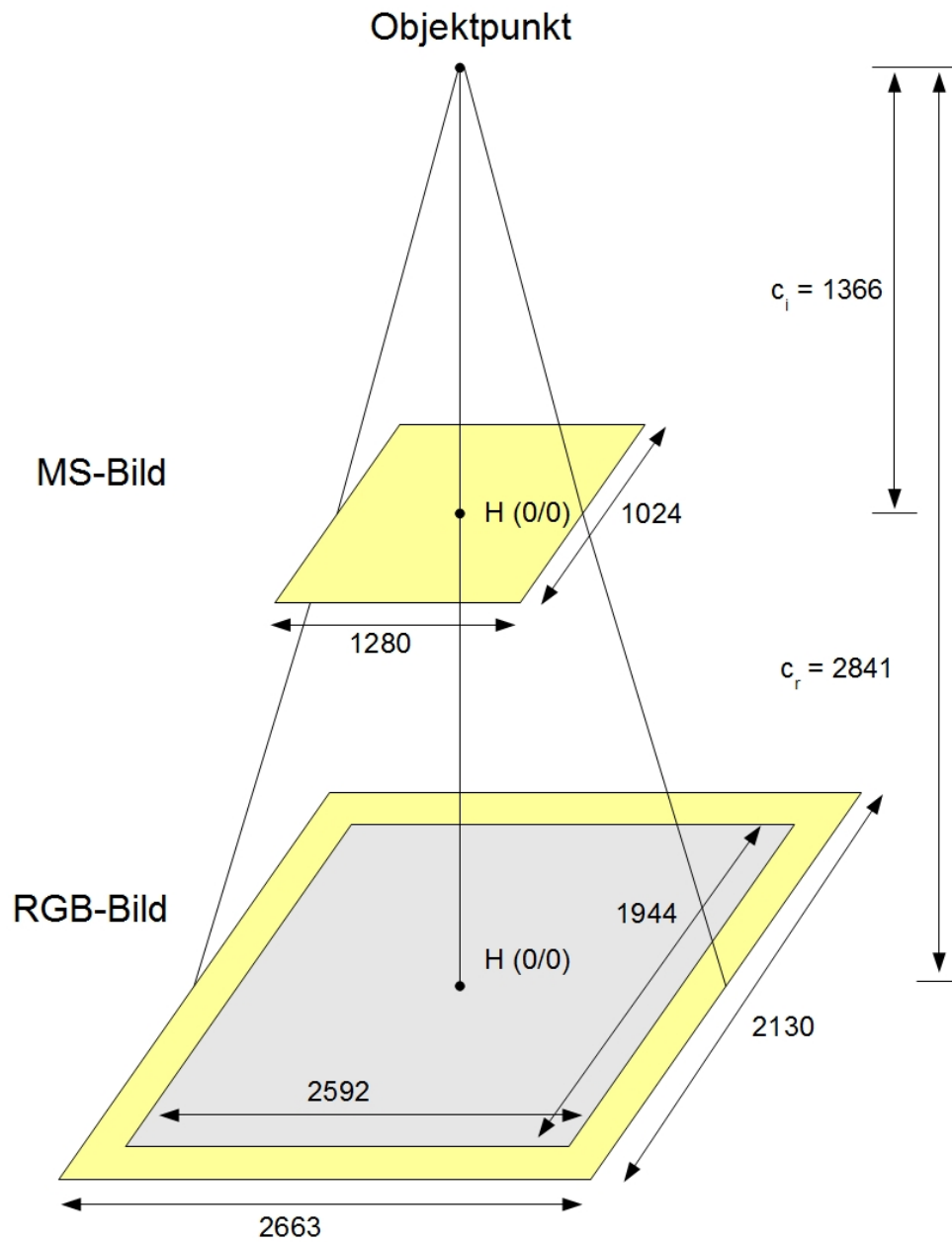


Abbildung 3.3: Anpassung der Kamerakonstante vom Multispektralbild an die Kamerakonstante des RGB-Bildes.

3 Verfahren zur Detektion von Blattkrankheiten

Mathematisch erfolgte die Anpassung der Kamerakonstante folgendermaßen. Für beide Kameras ist die Kalibriermatrix aufgestellt worden. Für die Multispektralkamera lautet sie

$$K_i = \begin{bmatrix} c_i & 0 & x_i \\ 0 & c_i & y_i \\ 0 & 0 & 1 \end{bmatrix}, \quad (3.1)$$

und für die RGB-Kamera

$$K_r = \begin{bmatrix} c_r & 0 & x_r \\ 0 & c_r & y_r \\ 0 & 0 & 1 \end{bmatrix}. \quad (3.2)$$

Die Umwandlung der Kamerakonstante kann dann folgendermaßen geschehen

$$K_{ir} = K_r K_i^{-1}. \quad (3.3)$$

Denn durch die Anwendung der inversen Kalibriermatrix K_i^{-1} auf einen Punkt $\mathbf{P}_i$ im Koordinatensystem der MS-Kamera wird dieser Punkt durch einen Richtungsvektor im Kamerasystem transformiert und durch die darauf folgende Multiplikation mit der Kalibriermatrix K_r in das Koordinatensystem der RGB-Kamera überführt.

Für K_{ir} ergibt sich

$$K_{ir} = \begin{bmatrix} \frac{c_r}{c_i} & 0 & -\frac{c_r x_i}{c_i} + x_r \\ 0 & \frac{c_r}{c_i} & -\frac{c_r y_i}{c_i} + y_r \\ 0 & 0 & 1 \end{bmatrix}. \quad (3.4)$$

Normalisierung der Multispektralbilder für die Zuordnung Nach Anpassung der Kamerakonstante können die Bilder beider Kameras gemeinsam mit der Software AURELO orientiert werden. Der erste Schritt für eine gemeinsame Orientierung besteht darin, in den Bildern identische Punkte zu identifizieren. Punkte werden dann als identisch bezeichnet, wenn sie denselben Objektpunkt repräsentieren. Dieses sogenannte Matching der Bilder wird bei der Software AURELO mittels des SIFT-Operators von Lowe durchgeführt (Läbe und Förstner, 2006). Dieser Operator ist invariant gegen Skalierung, Rotation und Translation. Zudem ist er relativ robust gegen Beleuchtungsänderungen und Bildrauschen (siehe Lowe, 2004). Dies ist sehr vorteilhaft, da

die MS-Bilder einerseits ein deutlich höheres Rauschen und andererseits eine deutlich geringere Helligkeit als die RGB-Bilder aufweisen.

Die Entwicklung des SIFT-Operators erfolgte allerdings anhand von Grauwertbildern, welche auf den Farbinformationen von RGB-Bildern basierten. Dadurch liefert der Operator nur dann wirklich gute Ergebnisse, wenn die Grauwertverteilung in den verschiedenen Bildern einer Szene relativ einheitlich ist. Im vorliegenden Fall sollen aber identische Punkte in MS- und RGB-Bildern gefunden werden, d.h. in Bildern mit Farbinformationen unterschiedlicher Kanäle und damit unterschiedlicher Grauwertverteilungen.

Da der Operator auf Grauwertbildern arbeitet, werden die Farbbilder normalerweise anhand folgender Standardformel in ein Grauwertbild umgerechnet

$$I_{rgb} = 0.299 R + 0.587 G + 0.114 B . \quad (3.5)$$

In dieser Gleichung steht R für den Rotkanal, G für den Grünkanal und B für den Blaukanal. Der Wert I_{rgb} gibt die Luminanz an (siehe Poynton, 2003).

Da das MS-Bild keinen Blaukanal aufweist, kann diese Standardformel nicht zur Umrechnung des MS-Bildes in ein Grauwertbild angewandt werden. Wie in Abbildung 3.4 zu sehen ist, werden aufgrund der unterschiedlichen Grauwerte der RGB- und MS-Bilder nur sehr wenige, völlig falsche identische Punkte gefunden. Es muss daher ein anderer Weg gefunden werden, Grauwertbilder aus den RGB- und MS-Bildern zu erzeugen.

Ein Weg wäre, die Farbinformation nur von einem Kanal zu wählen. Wie bereits in der Beschreibung des Kamerasystems in Kapitel 3.1.2 eingeführt, zeichnen sowohl die MS- als auch die RGB-Kamera den Rot- und den Grünkanal auf. Betrachtet man allerdings die Rot- (siehe Abbildung 3.5) und Grünkanäle (siehe Abbildung 3.6) näher, stellt man fest, dass die aufgezeichneten Wellenlängenbereiche vor allen Dingen im Grünkanal voneinander abweichen müssen.

Aus diesem Grunde wurde experimentell ein Verfahren entwickelt die MS-Bilder mittels Bildverarbeitungsmethoden (siehe Jähne, 2005) zu normalisieren und damit an die Grauwerte der RGB-Bilder anzupassen:

1. RGB-Bild anhand der Standardformel (3.5) in ein Graubild umwandeln.
2. Histogramm des Grauwertbildes I_{rgb} nach folgender Formel erstellen

$$h(g) = \#\text{Pixel mit Wert } g .$$

3. Multispektralbild anhand folgender Formel in ein Grauwertbild umwandeln

$$I_{ms} = 0.2 I + 0.75 R + 0.05 G ,$$

wobei I für den Infrarotkanal steht. Die Bestimmung der in dieser Formel angegebenen Gewichte erfolgte experimentell.

3 Verfahren zur Detektion von Blattkrankheiten

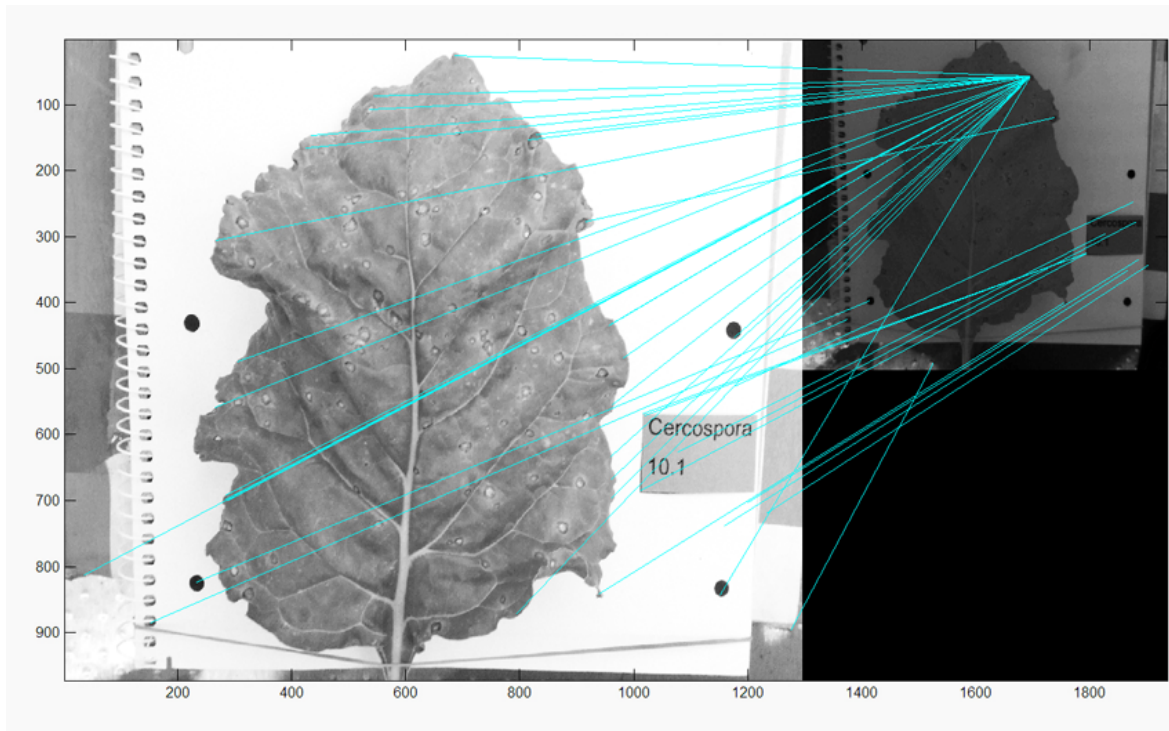


Abbildung 3.4: Low Matching zwischen Graubildern eines RGB- und Multispektralbildes. Beide Bilder wurden mit Hilfe der Standardformel in ein Graubild umgerechnet und dann das Lowe-Matching darauf ausgeführt. Es wurden nur sehr wenige, völlig falsche Matches gefunden.

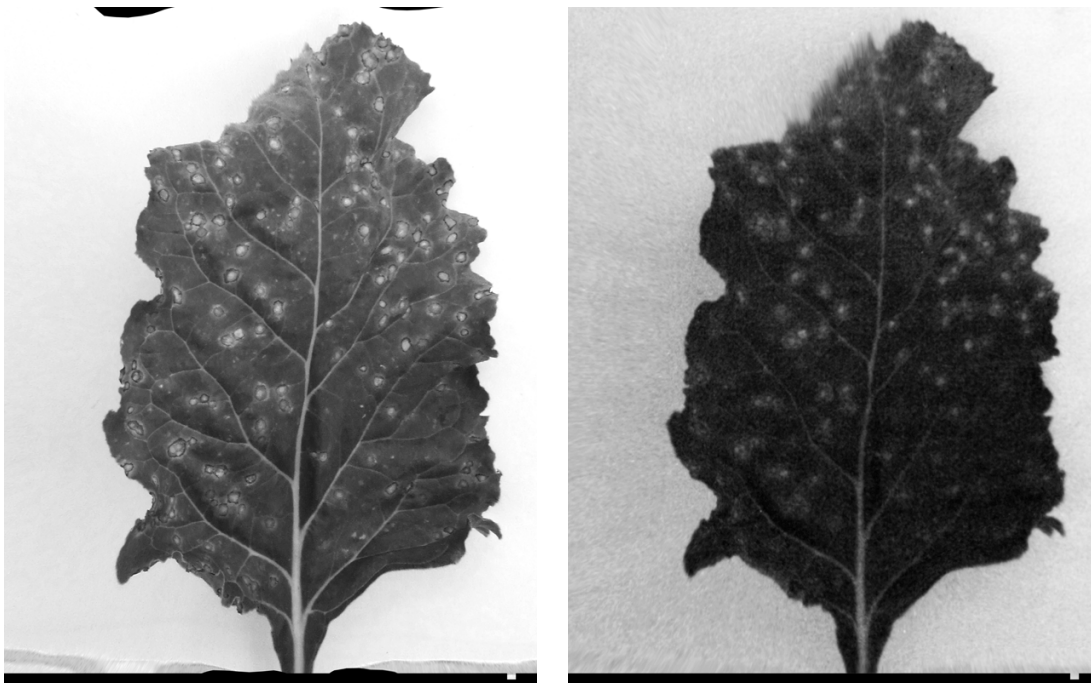


Abbildung 3.5: Rotkanal links des RGB-Bildes und rechts des Multispektralbildes.

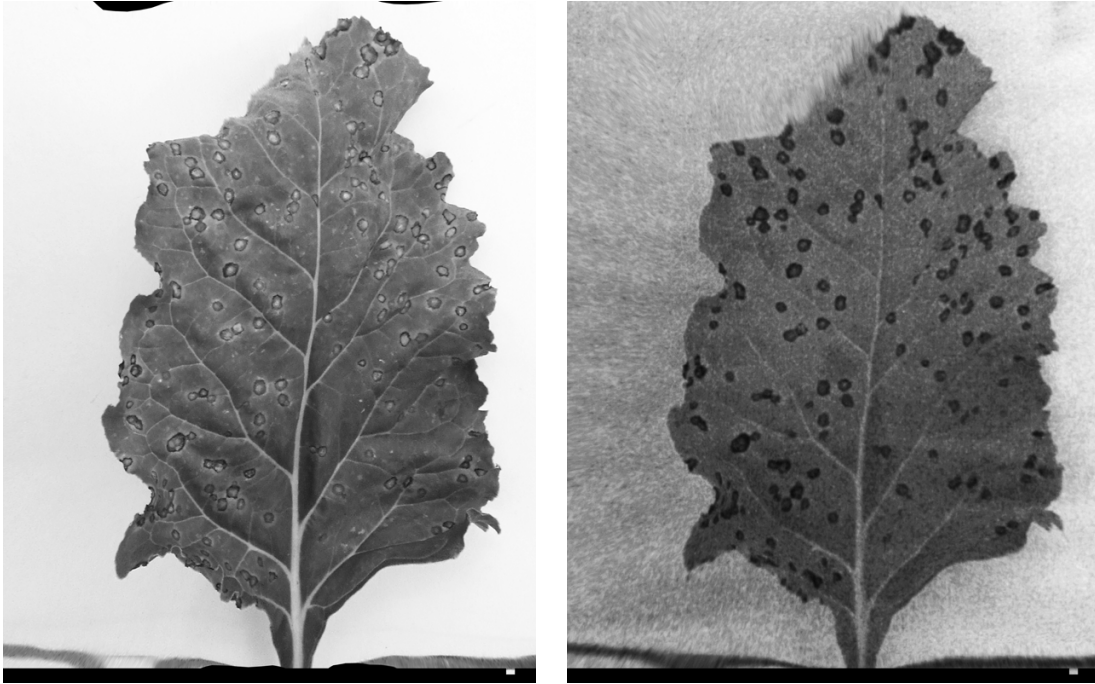


Abbildung 3.6: Links der Grüncanal des RGB-Bildes und rechts der Grüncanal des Multispektralbildes.

4. Histogramm des MS-Grauwertbildes I_{ms} an das Histogramm des RGB-Grauwertbildes $h(g)$ angleichen.

Wendet man den SIFT-Operator auf den so umgewandelten Grauwertbildern an, erhält man deutlich mehr Matches mit einer deutlich höheren Qualität (siehe Abbildung 3.7).

3.2.2 Erstellung eines multispektralen Orthobildes

Als Ergebnis der zuvor beschriebenen gemeinsamen Orientierung der Farb- und Multispektralaufnahmen liegen die Projektionsmatrizen der einzelnen Bilder vor. Anhand dieser Projektionsmatrizen kann ein Oberflächenmodell des jeweiligen Blattes erstellt werden. Dieses Modell wurde mit der INPHO-Software MATCH-T (Lemaire, 2008) erstellt. In diesem Oberflächenmodell existiert zu jedem Oberflächenpunkt P_{xyz} genau eine Tiefeninformation z . Damit handelt es sich genau genommen um 2 1/2-D Modelle, denn in einem vollständigen 3D-Modell können zu jedem Punkt P beliebig viele Tiefeninformationen z existieren. Für die hier durchgeführte Sensorfusion reicht die 2 1/2-D Information aber völlig aus. Der Einfachheit halber wird in der ganzen Arbeit von 3D-Modellen gesprochen, obwohl es sich also korrekterweise um 2 1/2 D-Modelle handelt.

Wie in Abbildung 3.2 auf Seite 32 zu sehen ist, sind zu jedem Oberflächenpunkt P_{xyz} die Farbinformationen aus den einzelnen Bildern bekannt. Der Übersicht halber

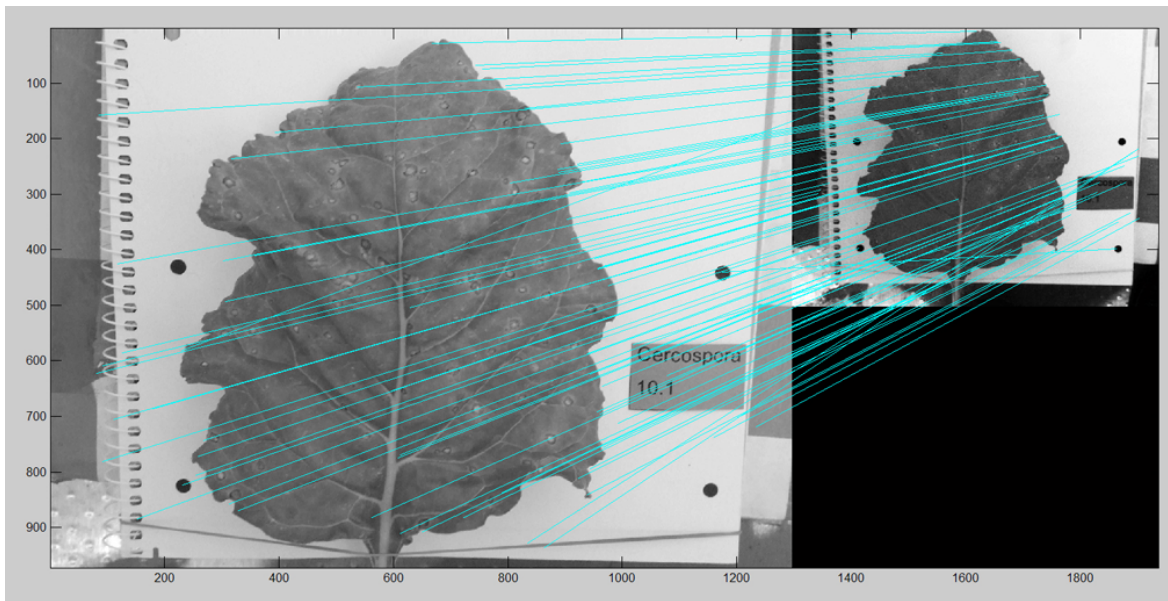


Abbildung 3.7: Lowe Matching zwischen Graubildern eines RGB- und Multispektralbildes. Das RGB-Bild wurde mit Hilfe der Standardformel in ein Graubild umgerechnet und das Multispektralbild mit der abgewandelten Formel und einer Anpassung des Histogramms des MS-Grauwertbildes an das Histogramm des RGB-Grauwertbildes. Es wurden doppelt so viele Matches als im vorigen Fall gefunden, wobei die Qualität der Matches erheblich höher ist.

sind in dieser Abbildung nur zwei RGB-Bilder und ein MS-Bild eingezeichnet. In der Realität wurden aber, wie in Kapitel 3.1.2 eingeführt, jeweils vier RGB-Bilder erstellt. Dies führt dazu, dass zu jedem Oberflächenpunkt P_{xyz} folgende 15 Farbwerte bekannt sind:

- $4 \times$ blau aus den vier RGB-Bildern,
- $5 \times$ grün aus den vier RGB-Bildern und dem MS-Bild,
- $5 \times$ rot aus den vier RGB-Bildern und dem MS-Bild,
- $1 \times$ infrarot aus dem MS-Bild.

Die 3D-Oberflächenpunkte P_{xyz} wurden dann auf eine darunterliegende Referenzebene projiziert, so dass ein rektifiziertes 2D-Bild entstand. In diesem sogenannten Orthofoto wurden die 3D-Oberflächenpunkte P_{xyz} auf die 2D-Punkte p_{xy} abgebildet. Die Erstellung des Orthofotos wurde mit dem Orthofotomodul „OrthoMaster“ der INPHO-Software (INPHO, 2008) erzeugt.

Die geometrische Genauigkeit der so generierten Orthofotos wurde manuell überprüft (siehe Kapitel 4.1). Dazu wurde kontrolliert inwieweit bei den unten dargestellten Ebenen in Abbildung 3.2 die Lage der Blattflecken und Blattränder übereinstimmt bzw. ob Verschiebungen aufgetreten sind. Dabei zeigte sich, dass bei 27% der Bilddatensätze die Ebenen eines RGB-Bildes und eines MS-Bildes übereinstimmen. Bei 56% der Datensätze zeigte sich stellenweise eine Verschiebung von bis zu 5 Pixeln und in 17% der Fälle war eine stärkere Verschiebung von bis zu 10 Pixeln aufgetreten. Diese Ungenauigkeiten können auf viele Ursachen zurückzuführen sein. So könnte z.B. bei der Bestimmung der identischen Punkte eine gewisse Abweichung aufgetreten sein oder bei der Erstellung der dichten 3D-Oberfläche sind durch die Interpolation Ungenauigkeiten hinzugekommen. Um die Ursachen explizit ausfindig zu machen, müssten weitergehende Untersuchungen durchgeführt werden. Um einen Einfluß der Verschiebungen auf die im folgenden durchgeführten Klassifikationen zu verhindern, wurden nur Datensätze ohne Verschiebung zur Klassifikation eingesetzt (siehe Kapitel 4.1).

3.3 Hierarchischer Klassifikationsprozess

Dieses Kapitel erläutert das Verfahren zur Detektion von Blattkrankheiten anhand der fusionierten, multispektralen Orthofotos. Das erste Unterkapitel gibt eine Übersicht über das gesamte Verfahren. Anschließend werden die Test- und Referenzdaten vorgestellt sowie die Bewertungsstrategie dargelegt, anhand dessen die Qualität der Klassifikationsergebnisse beurteilt wurde. Im Anschluß daran werden die einzelnen Schritte des Klassifikationsverfahrens näher beschrieben.

3.3.1 Übersicht über das Gesamtverfahren

Die Abbildung 3.8 gibt einen Überblick über das entwickelte Gesamtverfahren zur Detektion von Blattkrankheiten. Wie man in dieser Abbildung sehen kann, erfolgt im ersten Schritt eine pixelweise, adaptive Bayesklassifikation. Die Bayesklassifikation basiert auf den Grundlagen in Kapitel 2.3 und wurde entsprechend des folgenden Unterkapitels 3.3.4 zur Klassifikation von Blattkrankheiten angepasst. In diesem Schritt wurde die pixelweise Klassifikation einer regionenbasierten Klassifikation vorgezogen, da die zur regionenbasierten Klassifikation notwendige Segmentierung des Blattes, z.B. mit Hilfe eines Wasserscheidenalgorithmus', keinen Erfolg brachte (siehe Kapitel 4.4.1).

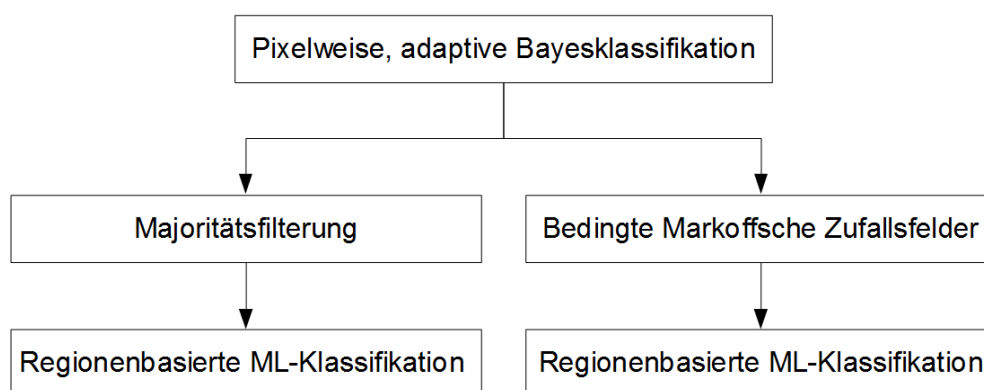


Abbildung 3.8: Die Grafik zeigt die Übersicht über das Gesamtverfahren. Im ersten Schritt erfolgt in jedem Fall eine pixelweise, adaptive Bayesklassifikation. Das daraus resultierende Labelbild wird entweder durch eine Majoritätsfilterung oder durch Anwendung von bedingten Markoffschen Zufallsfeldern geglättet. Die durch die Glättung entstandenen Regionen werden dann im letzten Schritt noch durch eine ML-Klassifikation nachklassifiziert. Die Vor- und Nachteile der beiden Glättungsarten wird in Kapitel 4.3 diskutiert.

Im zweiten Schritt erfolgt eine Glättung des pixelweisen Klassifikationsergebnisses. Die Glättung kann dabei entweder durch eine Majoritätsfilterung erfolgen oder durch den Einsatz von bedingten Markoffschen Zufallsfeldern. Die Grundlagen der bedingten Markoffschen Zufallsfelder können den Arbeiten von Kumar *et al.* (2005) und Korč und Förstner (2008) entnommen werden. Welche Vor- und Nachteile die beiden Glättungsverfahren bieten, wird im Kapitel 4.3 untersucht.

Im letzten Schritt erfolgt schließlich eine ML-Klassifikation der durch die pixelweise Klassifikation und anschließenden Glättung segmentierten Regionen. Diese Regionenklassifikation ist im Unterkapitel 3.3.6 näher beschrieben. Sie dient hauptsächlich dazu, fälschlicherweise als Blattflecken klassifizierte Regionen wieder der Klasse „Gesund“ zu-

zuführen und damit die Fehlklassifikationsrate dieser Klasse zu reduzieren.

3.3.2 Test- und Referenzdaten

Als Testdaten für die Entwicklung des Klassifikationsverfahrens zur Detektion von Blattkrankheiten dienten die fusionierten Multispektral-Orthofotos. Die Überprüfung der geometrischen Genauigkeit ergab allerdings, dass die Bilder teilweise Verschiebungen zwischen den RGB und Infrarotebenen aufweisen. Diese Verschiebungen traten häufig nicht auf dem gesamten Blatt auf, sondern nur stellenweise. Waren von den Verschiebungen auch Blattflecken betroffen, wurden die jeweiligen Bilddatensätze komplett aus dem Testdatensatz entfernt (siehe Kapitel 4.1). Letztendlich befanden sich damit in dem Testdatensatz 98 *Cercospora*- und 145 *Uromyces*-Bilder in allen Entwicklungsstadien der beiden Krankheiten.

Zu diesen Bildern wurden Referenzdaten durch manuelle Klassifikation der Bilder erzeugt. Die manuelle Klassifikation erfolgte mit dem sogenannten Annotation Tool (Korč und Schneider, 2007), welches für die Annotation von Blättern mit Blattkrankheiten angepasst wurde. Bei der Annotation wurden die Blattränder und Blattflecken mit Hilfe der Maus umfahren und den jeweiligen Klassen 'Blatt', '*Cercospora beticola*' oder '*Uromyces betae*' zugeordnet. Die so erstellten Referenzdaten wurden als korrekt angenommen und sowohl zum Training der Klassifikatoren als auch zum Bewerten der Klassifikationsergebnisse verwendet.

Als Klassifikationsmerkmale stehen, wie in Kapitel 3.2.2 erläutert, prinzipiell zu jedem Pixel 15 Farbwerte zur Verfügung. Zwölf dieser Farbwerte beruhen allerdings auf den vier Aufnahmen mit ein und derselben Kamera. Da die Aufnahmen unter einheitlicher Beleuchtung aufgenommen wurden, ist durch diese doppelte Information keine Verbesserung in der Klassenseparabilität zu erwarten. Desweiteren brachte in Untersuchungen auch der Rot- und Grünkanal der MS-Kamera keine Erhöhung der Klassifikationsgenauigkeiten. Aus diesem Grunde basieren die Bilder im Testdatensatz pro Pixel auf den Kanälen Rot, Grün und Blau von einem der RGB-Bilder und zusätzlich auf der Infrarotinformation des MS-Bildes.

In einem ersten Experiment zur Abschätzung der erreichbaren Klassifikationsgenauigkeiten bei der Detektion von Blattkrankheiten mit dem k NN-Klassifikator (siehe Kapitel 2.2) wurde zudem untersucht, inwieweit die Erweiterung des Merkmalsvektors um die Farbinformation der 4-er Nachbarschaft das Klassifikationsergebnis verbessert. Wie in Abbildung 3.9 grafisch illustriert, steht links für den Klassifikator nur die Farbinformation Rot, Grün, Blau und Infrarot des jeweiligen Pixels zur Verfügung. Im rechten Fall wird der Merkmalsvektor um die Farbinformation des oberen, rechten, linken und unteren Nachbarn erweitert, so dass hier insgesamt pro Pixel 20 Farbwerte zur Verfügung stehen. Das Ergebnis des Experimentes zeigte (siehe Kapitel 4.2.2), dass sich die Klassifikationsgenauigkeiten durch die Erweiterung des Merkmalsvektors beträchtlich verbesserten. So konnte z.B. die Klassifikationsrate des Braunrostes von

59% auf 76% erhöht werden. Das im folgenden vorgestellte Verfahren verwendet daher nur Bilder mit dem um die 4-er Nachbarschaft erweiterten Merkmalsvektor.

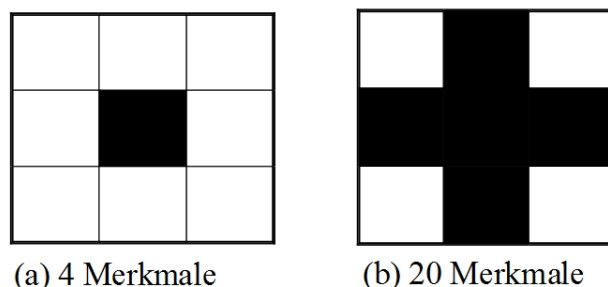


Abbildung 3.9: Pro Pixel gibt es die Merkmale Rot, Grün, Blau und Infrarot. Bei dem erweiterten Merkmalsvektor kommen dazu noch die Farbinformationen der 4-er Nachbarschaft, d.h. des rechten, linken, oberen und unteren Nachbarn. Links: Ohne Nachbarschaftsinformation, nur die Farbinformation des jeweiligen Pixels geht in den Merkmalsvektor ein. Rechts: Mit Nachbarschaftsinformation, die Farbinformationen der 4 Nachbarn wird mit berücksichtigt (siehe Bauer *et al.*, 2011).

3.3.3 Bewertungsstrategie bei der Klassifikation von Blattkrankheiten

Zur Entwicklung eines geeigneten Klassifikationsverfahrens ist es unabdingbar notwendig, entscheiden zu können, welches Klassifikationsergebnis als gut und welches als schlecht zu bewerten ist. Dies hängt unmittelbar damit zusammen, welches Ziel bei der Klassifikation erreicht werden soll. So wird häufig die Qualität eines Klassifikators anhand der Gesamtklassifikationsgenauigkeit bestimmt, d.h. es wird die Anzahl richtig klassifizierter Pixel zur Gesamtpixelanzahl in Beziehung gesetzt, unabhängig davon zu welcher speziellen Klasse die Pixel gehören. Dieses Vorgehen ist aber bei Vorhandensein von Klassen mit einer geringen Auftrittswahrscheinlichkeit eher ungünstig. So hat z.B. die Klasse Braunrost eine *a-priori* Wahrscheinlichkeit von 0.06% und die Klasse Gesund eine *a-priori* Wahrscheinlichkeit von 99.94%. Würde man nun einfach alle Braunrostpixel als Gesund fehlklassifizieren, würde die Gesamtklassifikationsgenauigkeit immer noch sehr gute 99.94% betragen. Die Klassifikationsgenauigkeit von der Klasse Braunrost läge dagegen bei 0%. Aus diesem Grunde werden in dieser Arbeit immer nur die Klassifikationsgenauigkeit von einzelnen Klassen, wie in Gleichung 2.10 angegeben, und nie die Gesamtklassifikationsgenauigkeit betrachtet.

In der Regel ist es nun so, dass eine Erhöhung der Klassifikationsgenauigkeit der einen Klasse automatisch zu einer geringeren Klassifikationsgenauigkeit der anderen Klasse führt. Im Bereich der Präzisionslandwirtschaft sollte allerdings sowohl die Klassifikationsgenauigkeit der kranken Klassen (Sensitivität) als auch die Klassifikationsgenauigkeit der gesunden Bereiche (Spezifität) hoch sein. Denn wenn kranke Pflanzen

übersehen werden und daher nicht behandelt werden, breitet sich die Krankheit evt. doch über das ganze Feld aus und führt zu Ernteausschlag. Werden dagegen zu viele gesunde Pflanzen als krank deklariert und daher mitbehandelt, dürfte es zwar zu keinem Ernteausschlag führen, dafür wird aber mehr Spritzmittel ausgebracht als benötigt und damit das Ziel der Präzisionslandwirtschaft verfehlt. Das Ziel in dieser Arbeit ist daher für alle Klassen gleichermaßen eine möglichst hohe Klassifikationsgenauigkeit zu erreichen.

Desweiteren kann die Bewertung des Klassifikationsergebnisses, genauso wie die Klassifikation selbst, auf verschiedenen Ebenen geschehen. Und zwar auf der

- Pixel-,
- Regionen- oder
- Blattebene.

Auf der Pixelebene wird für jedes individuelle Pixel entschieden, ob die Klassenzuordnung korrekt vorgenommen wurde oder nicht.

Auf der Regionenebene wird die Bewertung der Klassenzuordnung dagegen für komplette Regionen, deren Pixel alle zur gleichen Klasse l zugewiesen wurden, durchgeführt. Zu diesem Zweck muss zunächst für jede segmentierte Region entschieden werden, zu welcher Klasse sie korrekter Weise zugeordnet werden müsste. In der Regel stimmen nämlich die segmentierten Regionen und die Regionen des Referenzbildes nicht exakt überein. Eine Idee ist es, die segmentierte Region derjenigen Referenzklasse zuzuordnen, zu der die Mehrzahl der Pixel dieser Region im Referenzbild gehören. Problematisch ist dieses Vorgehen allerdings bei Klassen, welche im Referenzbild nur sehr kleine Regionen aufweisen, wie z.B. dem Braunrost. Denn es kann z.B. passieren, dass die segmentierte Region deutlich größer ist als die kleine, korrekte Region. Dies würde dann dazu führen, dass die ganze segmentierte Region einer anderen Referenzklasse zugeordnet wird und die Klasse mit den kleinen Regionen völlig untergeht. Die Überprüfung des Klassifikationsergebnisses auf Regionenebene ist von daher also deutlich schwieriger und kann unterschiedlich ausfallen, je nachdem wie man die Zuordnung einer Region zu einer Ground Truth Klasse vornimmt. Im Experiment 4.4 sind die Ergebnisse der regionenbasierten Klassifikation daher ebenfalls auf der Pixelebene dargestellt worden.

Auf der Blattebene wird für das gesamte Blatt entschieden ob es gesund oder krank ist bzw. welche Krankheit es aufweist. Diese Entscheidung kann z.B. anhand des Anteils der kranken im Verhältnis zur gesunden Blattfläche oder der Anzahl der kranken Regionen getroffen werden. Die Entscheidung ab wann ein Blatt als krank deklariert werden soll, sollte aber von einem Phytomediziner definiert und dann entsprechend im System implementiert werden. Eine Bewertung auf Blattebene erfolgt daher in dieser Arbeit nicht.

3.3.4 Pixelweise, adaptive Bayesklassifikation

Wie in der Übersicht über das Gesamtverfahren in Kapitel 3.3.1 erläutert, besteht der erste Schritt in einer pixelweisen, adaptiven Bayesklassifikation (siehe Bauer *et al.*, 2011). Diese Klassifikation basiert auf der Bayesklassifikation mit Risikominimierung von Fukunaga (1972) (siehe Kapitel 2.3.1). Ein Pixel mit dem Merkmalsvektor $\mathbf{h}_s(\mathbf{y})$ wird demnach der Klasse $\hat{x}_s$ anhand der Formel 2.5 zugewiesen. Der Hauptunterschied zwischen dem hier im Folgenden vorgestellten adaptiven Bayesklassifikator und dem Bayesklassifikator von Fukunaga (1972) liegt in der Aufstellung der Kostenfunktion k . Denn bei der adaptiven Bayesklassifikation werden die Kosten für die verschiedenen Fehlentscheidungen nicht manuell bestimmt und entsprechend angegeben, sondern automatisch anhand eines iterativen Optimierungsprozesses zur Erhöhung der Klassifikationsgenauigkeiten ermittelt. Die resultierenden Werte sind daher im engeren Sinne keine Kosten, sondern eher Gewichtungen der einzelnen Klassen. Im Folgenden wird daher nicht mehr von einer Kostenfunktion, sondern nur noch von der Gewichtungsfunktion w gesprochen.

Zweck der Gewichtungsfunktion ist es, den selteneren Klassen *Cercospora beticola* und vor allen Dingen *Uromyces betae* ein höheres Gewicht zu geben und damit ihre Klassifikationsgenauigkeit zu erhöhen, ohne gleichzeitig die Klassifikationsgenauigkeit der Klasse Gesund zu stark zu reduzieren. In einem Experiment zur Bestimmung der *a-priori* Wahrscheinlichkeit $P(l)$ (siehe Kapitel 4.2.5) zeigte sich, dass die Klassifikationsgenauigkeiten der beiden Blattkrankheiten bei einer ML-Klassifikation deutlich höher lagen als bei einer MAP-Klassifikation. Aus diesem Grunde wurde zur Bestimmung der Gewichtungsfunktion w zunächst eine Basis-Gewichtungsfunktion w_{basis} aufgestellt. Die Gewichte in dieser Basis-Gewichtungsfunktion wurden folgendermaßen abhängig von den *a-priori* Wahrscheinlichkeiten $P(l)$ der einzelnen Klassen l bestimmt

$$w_{\text{basis}}(l, x_s) = \begin{cases} \frac{1}{P(l)} & \text{falls } l \neq x_s \\ 0 & \text{falls } l = x_s \end{cases}. \quad (3.6)$$

Zur Berechnung dieser Basis-Gewichtungsfunktion w_{basis} erfolgte die Berechnung der dazu notwendigen *a-priori* Wahrscheinlichkeiten $P(l)$ durch eine Vorklassifikation der Trainingsdaten mit einem ML-Klassifikator. Im Vergleich zur Verwendung der Ground Truth *a-priori* Wahrscheinlichkeiten brachte dies um ca. 5% höhere Klassifikationengenauigkeiten bei der Klasse Gesund und *Cercospora beticola* (siehe Kapitel 4.2.6).

Im nächsten Schritt wurde dann anhand der Basis-Gewichtungsfunktion w_{basis} die eigentliche Gewichtungsfunktion w bestimmt. Dazu wurden die Trainingsdaten zunächst mit dem Bayesklassifikator unter Verwendung Basis-Gewichtungsfunktion w_{basis} klassifiziert und die Gewichtungsfunktion w als Null-Matrix initialisiert. Das Klassifikationsergebnis wurde in der Konfusionsmatrix o abgespeichert und die Gewichtungsfunktion w folgendermaßen iterativ adaptiert

$$w_{\nu+1}(l, x_s) = w_{\nu}(l, x_s) + w_{\text{basis}}(l, x_s) \cdot o(l, x_s) \quad (3.7)$$

Um die Klassifikationsgenauigkeit aller drei Klassen gegen 100% zu konvergieren, wurde die Klassifikation der Trainingsdaten unter Verwendung der Gewichtungsfunktion w iterativ wiederholt bis der Änderungsfaktor a kleiner als 0.01 blieb:

$$a = \sum_{l, x_s} |o(l, x_s) - o_{\text{old}}(l, x_s)|. \quad (3.8)$$

Im Trainingsprozess des Bayesklassifikators muss außerdem die Likelihood-Funktion $L(\mathbf{h}_s(\mathbf{y}) | l)$ bestimmt werden. Es wurde hier ein Gauß'sches Mischmodell angenommen und die optimale Anzahl der Verteilungen für jede Klasse experimentell bestimmt (siehe Kapitel 4.2.4). Folgende Anzahlen wurden dabei für die drei Klassen ermittelt:

- Gesund – 2 Verteilungen
- *Cercospora beticola* – 3 Verteilungen
- *Uromyces betae* – 1 Verteilung

Die Parameter des Gauß'schen Mischmodells wurden für jede Klasse anhand des EM-Algorithmus' von Bilmes (1998) berechnet (siehe Kapitel 2.3).

Normalerweise werden desweiteren die *a-priori* Wahrscheinlichkeiten $P(l)$ der einzelnen Klassen in der Trainingsphase berechnet. Da aber die Auftrittswahrscheinlichkeit der Blattkrankheiten einerseits von Blatt zu Blatt und andererseits abhängig vom Entwicklungsstadium der Krankheit sehr unterschiedlich ist, machen vorab bestimmte *a-priori* Wahrscheinlichkeiten bei Blattkrankheiten wenig Sinn. Die *a-priori* Wahrscheinlichkeiten $P(l)$ wurden daher nicht in der Trainingsphase ermittelt, sondern in der Testphase für jedes Bild durch eine Vorklassifikation mit einem ML-Klassifikator unter Verwendung der obigen Likelihood-Funktion $L(\mathbf{h}_s(\mathbf{y}) | l)$ individuell berechnet. Wie dem Experiment in Kapitel 4.2.5 entnommen werden kann, weichen die Klassifikationsgenauigkeiten bei dieser Vorgehensweise nur minimal von den Klassifikationsgenauigkeiten unter Verwendung der Ground Truth *a-priori* Wahrscheinlichkeiten des jeweiligen Bildes ab.

3.3.5 Glättung der pixelweisen Klassifikation zur Regionenbildung

Nach der adaptiven Bayesklassifikation folgt im nächsten Schritt eine Glättung dieses pixelweisen Klassifikationsergebnisses. Wie in Abbildung 3.10 zu sehen ist, treten bei einer pixelweisen Klassifikation typische Pixelfehler auf. D.h. einzelne Pixel werden einer anderen, falschen Klassen zugewiesen als die benachbarten Pixel. Um diese einzelnen Pixelfehler zu eliminieren, soll die pixelweise Klassifikation geglättet werden. Dazu haben wir zwei Varianten untersucht (siehe Kapitel 4.3):

3 Verfahren zur Detektion von Blattkrankheiten

- die Glättung mit einem Majoritätsfilter und
- der Einsatz von bedingten Markoffschen Zufallsfeldern.

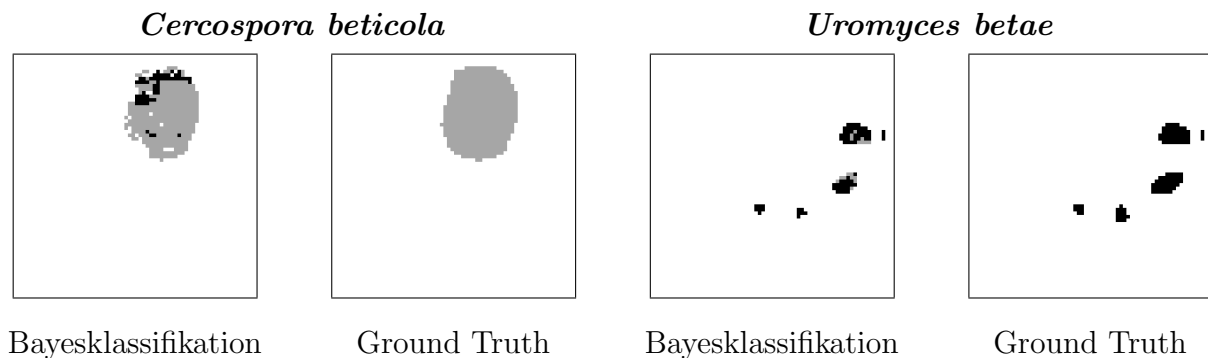


Abbildung 3.10: Ergebnisbilder der pixelweisen, adaptiven Bayesklassifikation. In den Bildern repräsentiert die weiße Farbe gesunde Blattfläche, grau die Klasse *Cercospora beticola* und schwarz die Klasse *Uromyces betae*. Die ersten beiden Ausschnitte zeigen das Klassifikationsergebnis und das dazugehörige Ground Truth-Bild von einem mit *Cercospora beticola* infizierten Blatt. Die letzten beiden Ausschnitte zeigen analog das Ergebnis eines mit *Uromyces betae* infizierten Blattes. Die Ausschnitte zeigen gut die einzelnen Pixelfehler und Fehlklassifikationen zwischen den beiden Blattkrankheiten (siehe Bauer *et al.*, 2011).

Bei der Anwendung einer Majoritätsfilterung (siehe Oommen *et al.*, 2004) wird die Klasse des jeweiligen Pixels und der 4er-Nachbarschaft, d.h. des linken, rechten, oberen und unteren Nachbarn, betrachtet. Das jeweilige Pixel erhält dann diejenige Klasse, welche am häufigsten vertreten ist.

Die Anwendung der bedingten Markoffschen Zufallsfeldern, so wie wir sie hier verwenden, ist im Wesentlichen in Bauer *et al.* (2011) dargestellt. Lediglich die Werte in der Matrix T wurden in dieser Arbeit durch Experimente empirisch neu ermittelt. In dieser Matrix T wird angegeben, wie wahrscheinlich das direkt benachbarte Auftreten von zwei Klassen ist. Ist die Wahrscheinlichkeit für eine Klassenkombination $x_s, x_{s'}$ hoch, ist auch der Wert an der jeweiligen Stelle in der Matrix T hoch. Entsprechend ist der Wert in der Matrix bei einer geringen Wahrscheinlichkeit für eine bestimmte Klassenkombination niedrig. Die Reihenfolge der Klassen ist dabei irrelevant, d.h. es gilt $x_s, x_{s'} = x_{s'}, x_s$.

Folgende Werte haben in den hier durchgeführten Experimenten die besten Ergebnisse gebracht:

$$T = \begin{bmatrix} 0.3 & 0.01 & 0.02 \\ 0.01 & 0.4 & 0.02 \\ 0.02 & 0.02 & 0.02 \end{bmatrix}$$

Die erste Zeile und Spalte gibt die Werte für die Blattfleckenkrankheit *Cercospora beticola* an, die zweite Zeile und Spalte die Werte für den Braunrost *Uromyces betae* und die letzte Zeile und Spalte die Werte für die gesunde Blattfläche. Wie man sehen kann, erhielt der Braunrost eine besonders hohe Wahrscheinlichkeit, um zu verhindern, dass er aufgrund seiner Seltenheit einfach weggeglättet wird. Dafür musste zudem die Wahrscheinlichkeit der gesunden Blattfläche in dieser Matrix deutlich reduziert werden. Sie bekam somit die gleiche Wahrscheinlichkeit wie das benachbarte Auftreten von Blattkrankheit und gesunder Blattfläche. Die geringste Wahrscheinlichkeit erhielt das benachbarte Auftreten von den beiden Blattkrankheiten. Auf diese Weise konnte effektiv verhindert werden, dass eine Fehlklassifikation von Braunrost am Rand der *Cercospora*-Blattflecken auftritt (siehe Kapitel 4.3.2).

3.3.6 Regionenbasierte Maximum-Likelihood Klassifikation

Die im vorigen Schritt durchgeführte Glättung des pixelweisen Klassifikationsergebnisses führte zur Eliminierung der typischen Pixelfehler und vor allen Dingen bei Verwendung der bedingten Markoffschen Zufallsfeldern auch zu deutlich glatteren Rändern der Blattflecken (siehe Kapitel 4.3). Es existieren allerdings immer noch eine nicht unerhebliche Anzahl 'falscher' Blattflecken, welche eigentlich zur Klasse Gesund gehören. Dies ist z.B. in den Abbildungen 4.8, 4.9 und 4.10 auf den Seiten 76, 77 und 78 ersichtlich. Desweiteren besteht ein Blattfleck trotz Glättung des öfteren immer noch aus zwei Krankheiten. In der Regel sollte aber eine als krank erkannte Region auch nur einer Krankheit zugewiesen werden. Aus diesem Grunde erfolgt zur eindeutigen Klassifikation der Blattkrankheiten und zur Erhöhung der Klassifikationsgenauigkeit der Klasse Gesund eine regionenbasierte Klassifikation.

Der erste Schritt der regionenbasierten Klassifikation besteht darin, einheitliche Regionen zu bilden. Dazu wurde die gesunde Blattfläche als Hintergrund deklariert und alle als *Cercospora* oder *Uromyces* klassifizierten Pixel als Blattfleck. Zur Reduzierung der 'falschen' Blattflecken wurden dann die so entstandenen Regionen mittels eines ML-Klassifikators den Klassen Gesund oder Krank zugeordnet. Diese Klassifikation erfolgte anhand des mittleren Rot, Grün, Blau und Infrarotwerts der jeweiligen Region. Die Form der Blattflecken blieb in diesem Schritt unberücksichtigt, da die 'falschen' Blattflecken ganz unterschiedliche Größen und Formen aufweisen. Im Idealfall sollten auf diesem Wege alle 'falschen' Blattflecken eliminiert werden, so dass nur noch tatsächlich kranke Regionen auf den Blättern vorhanden sind.

3 Verfahren zur Detektion von Blattkrankheiten

Diese Regionen wurden nun in einem zweiten Schritt anhand der Fläche, des Umfangs und der Rundheit der Flecken der Blattfleckenkrankheit *Cercospora beticola* oder dem Braunrost *Uromyces betae* zugewiesen. In diesem Fall macht eine Klassifikation anhand der Größe und Form Sinn, da *Cercospora*-Flecken in der Regel erheblich größer sind als Rostpusteln und Rostpusteln zudem normalerweise kreisrund sind, während *Cercospora*-Flecken rund bis oval sind und sich bei Zusammenfließen mehrerer Flecken unregelmäßig ausbreiten (siehe Kapitel 1).

In beiden Schritten wurde eine Maximum-Likelihood Klassifikation durchgeführt. Es wurde dabei für alle Klassen jeweils eine Gaußsche Mischverteilung aus zwei Verteilungen angenommen. Die Parameter der Verteilungen wurden mit dem k -means Algorithmus vorab geschätzt und dann mit dem EM-Algorithmus (siehe Kapitel 2.3) optimiert.

Als Ergebnis der regionenbasierten Klassifikation konnte die Klassifikationsgenauigkeit der Klasse Gesund um bis zu 3% erhöht werden und die teilweise Fehlklassifikation der *Cercospora*-Blattflecken als Braunrost konnte stark reduziert werden. Die Klassifikationsgenauigkeit der seltenen Braunrostpusteln ist dagegen im Median erheblich von 88% auf 26% bei der Regionenklassifikation des Filterungsergebnisses bzw. von 83% auf 62% bei der Regionenklassifikation des Ergebnisses der bedingten Markoffschen Zufallsfelder abgesunken (siehe Kapitel 4.4).

4 Experimente

Die Entwicklung des im vorigen Kapitels vorgestellten Verfahrens zur Detektion von Blattkrankheiten basierte auf einigen Experimenten. Diese Experimente werden nun im Folgenden vorgestellt. Für jedes Experiment wird dabei der Zweck des Experiments erläutert, die Versuchsdurchführung dargelegt und die Ergebnisse des jeweiligen Experiments diskutiert.

Das erste Experiment diente zur empirischen Überprüfung der geometrischen Genauigkeit der Sensorfusion. Anhand der Datensätze, bei denen die Überprüfung positiv ausfiel, erfolgten im nächsten Schritt die Experimente zur Entwicklung der pixelweisen Klassifikationsstrategie. Daran schlossen sich die Experimente zur Glättung der pixelweisen Klassifikation und zur regionenbasierten Klassifikation basierend auf den geglätteten Bildern an. Im letzten Experiment wurde schließlich die Übertragbarkeit des Verfahrens auf andere Nutzpflanzen und Stressfaktoren untersucht.

4.1 Empirische Überprüfung der geometrischen Genauigkeit der Sensorfusion

Anhand des in Kapitel 3.2 vorgestellten Verfahrens wurden die Bilder der RGB- und MS-Kamera fusioniert. Zur Überprüfung der geometrischen Genauigkeit der Sensorfusion ist eine empirische Untersuchung durchgeführt worden, welche im Folgenden vorgestellt wird.

Die Überprüfung erfolgte anhand der in Kapitel 3.3.2 vorgestellten Annotation der Orthofotos. Wie dort beschrieben ist, wurde die Annotation jeweils auf dem RGB-Orthofoto durchgeführt. Zur Überprüfung wurde diese Annotation nun auf dem dazugehörigen MS-Orthofoto eingezeichnet. Anhand dessen wurde dann der Grad der Abweichung subjektiv eingeschätzt. Je nachdem wie exakt die Annotation mit dem Verlauf des Blattrandes und der Blattflecken auf dem MS-Bild übereinstimmte, wurde das Blatt folgenden Klassen zwischen 0 und 9 zugewiesen:

- 0 das Bild ist völlig ok.
- 1 das Bild hat keine Flecken und ist am Blattrand etwas verschoben.
- 2 das Bild hat keine Flecken und ist am Blattrand stärker verschoben.
- 3 das Bild hat Flecken, es ist aber nur der Blattrand etwas verschoben.

4 Experimente

- 4 das Bild hat Flecken, es ist aber nur der Blattrand stärker verschoben.
- 5 das Bild hat Flecken, sowohl der Blattrand als auch die Blattflecken sind etwas verschoben.
- 6 das Bild hat Flecken, sowohl der Blattrand als auch die Blattflecken sind stärker verschoben.
- 7 das Bild hat Flecken und ist am Blattrand etwas verschoben, ob die Flecken verschoben sind, ist nicht eindeutig zu erkennen.
- 8 das Bild hat Flecken und ist am Blattrand stärker verschoben, ob die Flecken verschoben sind, ist nicht eindeutig zu erkennen.
- 9 das Bild ist völlig ok, ob die Flecken verschoben sind, ist nicht eindeutig zu erkennen.

Sind der Blattrand oder die Blattflecken etwas verschoben, liegt eine Verschiebung von ungefähr bis zu 5 Pixeln vor. Unter einer stärkeren Verschiebung wurde eine Abweichung von ungefähr bis zu 10 Pixel verstanden. Dabei muss allerdings berücksichtigt werden, dass die Annotation selbst auch eine Toleranz von ungefähr bis zu 5 Pixeln aufweist.

Ein Bild, welches als völlig ok eingestuft wurde, ist in Abbildung 4.1 zu sehen. Eine geringfügige Verschiebung ist in Abbildung 4.2 dargestellt und Abbildung 4.3 zeigt zwei Ausschnitte von einem Blatt, bei dem der eine Ausschnitt völlig korrekt ist und im anderen Ausschnitt stärkere Abweichungen sichtbar sind.

Die manuelle Überprüfung wurde auf allen 180 vorhandenen *Cercospora*-Aufnahmen und auf allen 186 *Uromyces*-Aufnahmen durchgeführt. Bei den *Cercospora*-Aufnahmen wurden Klassenzuordnungen zwischen 0 und 7 vorgenommen; bei den *Uromyces*-Aufnahmen zusätzlich auch Zuordnungen in die Klassen 8 und 9, da hier eine Verschiebung der Flecken häufiger nicht eindeutig festgestellt werden konnte. Dies ist darauf zurückzuführen, dass die Rostpusteln erstens in der Regel nur wenige Pixel umfassen und zweitens die Rostpusteln durch das hohe Rauschen und die geringe Größe im MS-Bild nicht immer eindeutig zu identifizieren sind. Die Auftrittswahrscheinlichkeiten der einzelnen Klassen kann der Tabelle 4.1 entnommen werden.

Deklariert man sowohl die Klasse 0 als auch die Klasse 9 als völlig ok, hat die Fusionierung damit in 27.3% der Fälle einwandfrei funktioniert. Eine geringgradige Abweichung tritt in den Fällen 1, 3, 5 und 7 auf. Diese Klassen zusammen umfassen 55.5% der Blätter. Die restlichen Klassen 2, 4, 6 und 8 repräsentieren stärkere Abweichungen. Diese treten damit in insgesamt 17.2% der Fälle auf.

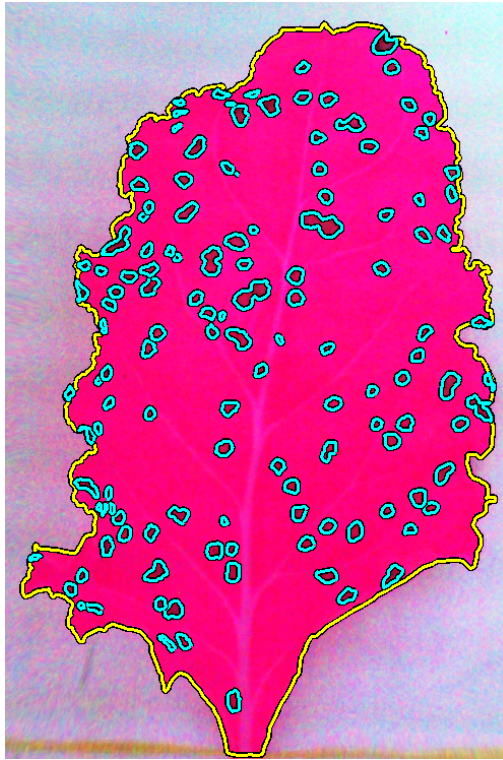


Abbildung 4.1: Korrektes Bild. Das Bild wurde dementsprechend der Klasse 0 zugeordnet.

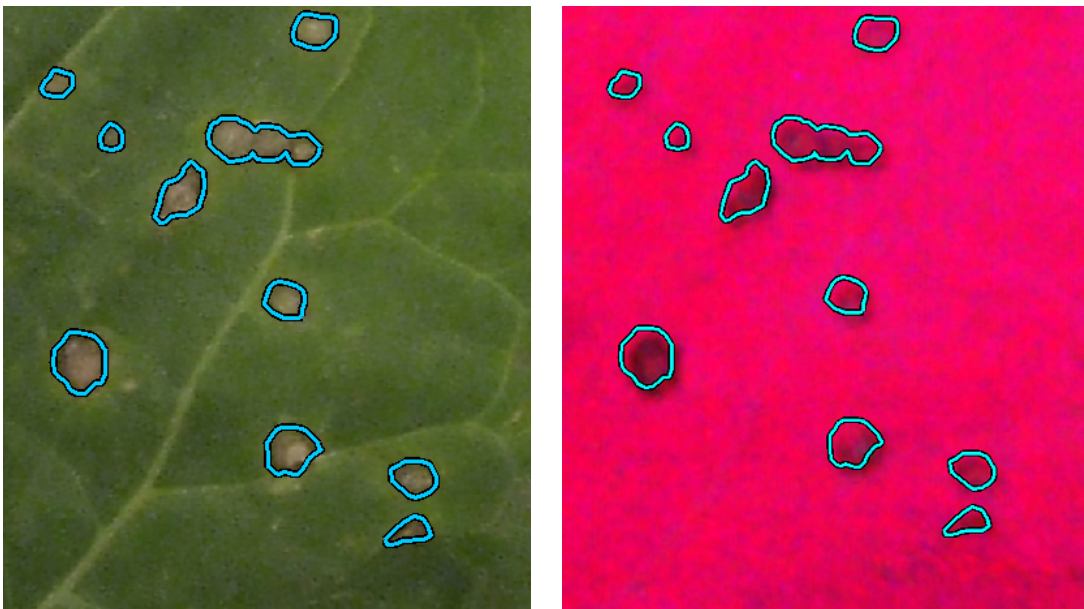


Abbildung 4.2: In diesem Bild ist die Geometrie größtenteils korrekt. Nur in dem abgebildeten Ausschnitt des Blattes sind die Flecken etwas verschoben.

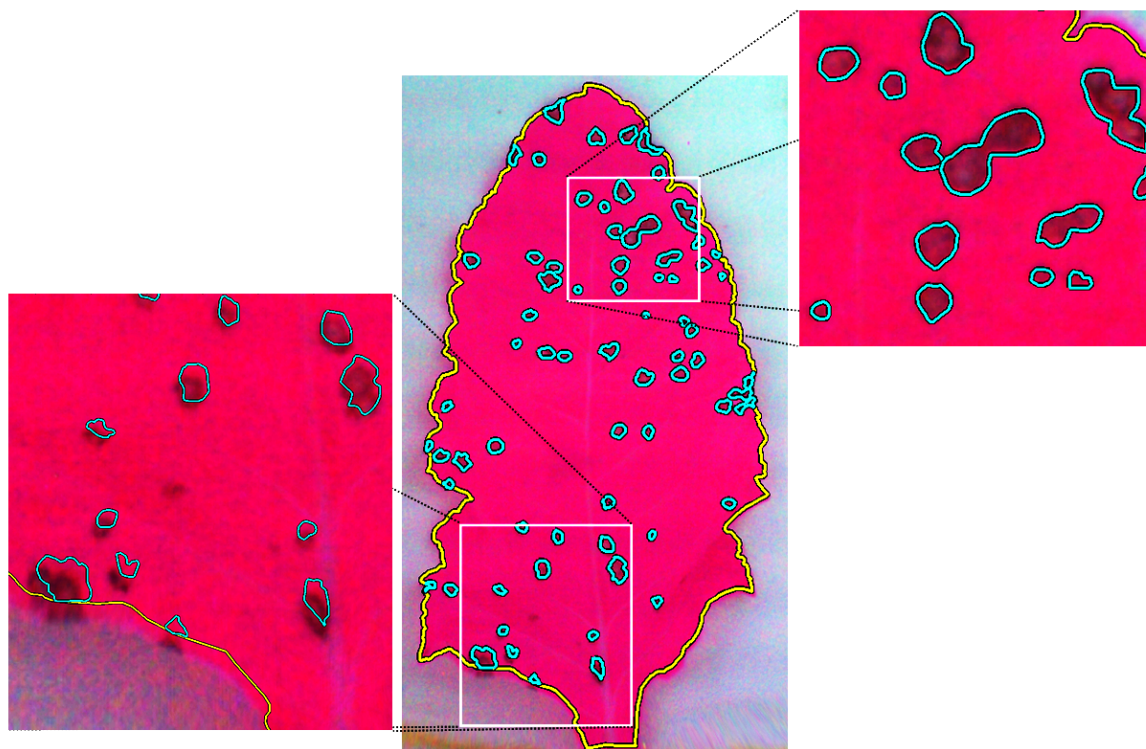


Abbildung 4.3: Auf diesem Bild sind der Blattrand und die Blattflecken teilweise stärker verschoben. In anderen Bildteilen aber auch völlig in Ordnung. In der Vergrößerung links unten ist eine stärkere Verschiebung sichtbar. In der Vergrößerung rechts oben stimmt dagegen alles komplett überein.

Geometrische Genauigkeiten	
Klasse	Auftrittswahrscheinlichkeit in %
0	24.3
1	24.0
2	7.6
3	6.0
4	0.3
5	14.8
6	8.2
7	10.7
8	1.1
9	3.0

Tabelle 4.1: Die Tabelle zeigt die Auftrittswahrscheinlichkeit in % der einzelnen Klassen basierend auf allen 366 Aufnahmen.

4.2 Entwicklung der pixelweisen Klassifikationsstrategie

Dieses Kapitel stellt die Experimente zur Entwicklung der pixelweisen Klassifikationsstrategie vor. Die pixelweise Klassifikation ist der erste Schritt zur Detektion der Blattkrankheiten (siehe Kapitel 3.3.1). Im Folgenden werden zunächst die in diesen Experimenten verwendeten Test- und Referenzdaten vorgestellt, welche auf den in Kapitel 3.3.2 vorgestellten Testbildern basieren.

Anschließend wird die Versuchsdurchführung mit dem k NN-Klassifikator beschrieben und das mit diesem Klassifikator durchgeführte Experiment zur Erweiterung des Merkmalsvektors um die 4-er Nachbarschaft vorgestellt. Das Ergebnis dieses Experimentes zeigt, dass eine Berücksichtigung der 4-er Nachbarschaft in den Merkmalen die Klassifikationsgenauigkeiten erheblich erhöht.

Daher basieren die folgenden Experimente zur Optimierung des adaptiven Bayesklassifikators alle auf dem erweiterten Merkmalsvektor. Zur Optimierung wurden folgende drei Experimente durchgeführt:

- Bestimmung der Anzahl der Gaußverteilungen pro Klasse,
- Bestimmung der *a-priori* Wahrscheinlichkeiten der einzelnen Klassen und
- Aufstellung der Gewichtungsfunktion für die Bayesklassifikation.

4.2.1 Verwendete Test- und Referenzdaten

Das Ziel der folgenden Experimente ist es, Klassifikatoren auf die bestmögliche Trennung zwischen kranker und gesunder Blattfläche hin zu optimieren. Dazu werden Bilder benötigt, welche Symptome der beiden Blattkrankheiten zeigen. Aus diesem Grund wurden für diese Experimente nur die Bilder ab dem zehnten Tag nach der Inokulation verwendet. Von dem in Kapitel 3.3.2 vorgestellten Test- und Referenzdatensatz erfüllen 55 *Cercospora*-Datensätze und 73 *Uromyces*-Datensätze diese Vorgabe.

Zur weiteren Minimierung der Rechenzeit wurden diese Bilder in Bildausschnitte der Größe 64×64 Pixel aufgeteilt. Die Bildausschnitte zeigten entweder

- nur Hintergrund,
- Hintergrund mit Blatt,
- nur gesunde Blattfläche oder
- sowohl gesunde als auch kranke Blattfläche.

Da, wie gesagt, für die folgenden Experimente nur die optimale Trennung zwischen gesunder und kranker Blattfläche relevant war, wurden nur diejenigen Bildausschnitte in den Testdatensatz aufgenommen, welche sowohl gesunde als auch kranke Blattfläche ohne Hintergrund zeigten. Der Testdatensatz umfasste damit 357 Bildausschnitte mit

Symptomen des Braunrostes *Uromyces betae* und 2256 Bildausschnitte mit Symptomen der Blattfleckenkrankheit *Cercospora beticola*. Aus diesen verbliebenen Bildausschnitten, welche also alle Krankheitssymptome in verschiedenen Stadien zeigten, wurden für jede Krankheit 350 Bildausschnitte zufällig ausgewählt.

Der im Folgenden eingesetzte Testdatensatz umfasste damit 700 Bildausschnitte der Größe 64×64 . Die Labelbilder mit den Ground Truth Klassen wurden analog zu den Daten-Bildern in Bildausschnitte aufgeteilt, so dass zu jedem Bildausschnitt ein passender Ausschnitt des Labelbildes existiert.

4.2.2 Berücksichtigung der 4-er Nachbarschaft in den Merkmalen

Das Ziel dieses Experimentes war es festzustellen, ob eine Erweiterung des Merkmalvektors um die Farbinformation der 4er-Nachbarschaft das Klassifikationsergebnis verbessert oder nicht. Dazu wurden zwei Experimente mit unterschiedlichem Merkmalsvektor durchgeführt. Wie in Abbildung 3.9 auf Seite 44 illustriert, umfasste der Merkmalsvektor im ersten Experiment (links in der Grafik) nur die vier Farbinformationen des jeweiligen Pixels selbst. Der Merkmalsvektor im zweiten Experiment (rechts in der Grafik) berücksichtigte dagegen auch die Farbinformationen des linken, rechten, oberen und unteren Nachbarn, so dass er insgesamt 20 Farbwerte beinhaltete. Zur Klassifikation wurde der in Kapitel 2.2 vorgestellte k -Nächste-Nachbarn (k NN) Klassifikator eingesetzt.

Versuchsdurchführung mit dem k NN Klassifikator

Wie in Kapitel 2.2 erläutert, wird zum Training des k NN-Klassifikators eine bestimmte Pixelanzahl mit den dazugehörigen Ground Truth Klassen x_s im Speicher abgelegt. Die Anzahl der Trainingsdaten ist dabei nicht ganz unwesentlich, denn für eine gute Abschätzung des Bayesfehlers sollte diese Anzahl möglichst groß sein. Je höher die Anzahl der Trainingsdaten aber ist, desto länger ist aufgrund des größeren Suchraumes auch die Berechnungszeit während der Klassifikation. Desweiteren sollte für jede Klasse möglichst die gleiche Anzahl an Trainingsdaten vorhanden sein.

Da in jedem Bildausschnitt im Durchschnitt nur zwölf *Uromyces*-Pixel enthalten sind, wurden für jede Klasse zufällig 2000 Pixel aus 350 der 700 Bildausschnitte zum Training ausgewählt. Die anderen 350 Bildausschnitte wurden dann zum Testen verwendet. Dieser Vorgang wurde, um eine hohe statistische Genauigkeit zu gewährleisten, fünf mal wiederholt.

Im Klassifikationsprozess wurde jeweils die Klassenzugehörigkeit der $k = 5$ nächsten Nachbarn betrachtet.

Ergebnisse des k NN Klassifikators

Die Ergebnisse der beiden Experimente mit und ohne Nachbarschaftsinformation in den Merkmalen sind in Tabelle 4.2 zu sehen. Jeweils der erste Wert pro Zeile und Zelle gibt das Ergebnis ohne Nachbarschaftsinformation an und entsprechend der zweite Wert das Ergebnis mit Nachbarschaftsinformation. Auf der linken Seite der Tabelle sind die Ground Truth Klassen angegeben. Jede Ground Truth Klasse ist durch eine Tilde, wie z.B. $\tilde{h}$, als solche markiert. Oben in der Tabelle sind die Testklassen aufgetragen. Sie enthalten die geschätzten Werte und werden durch das Dach, $\hat{h}$, repräsentiert. Die Klassenabkürzung h steht für die Klasse Gesund und die Abkürzung i für die verschiedenen Pathogene. Das spezielle Pathogen kann jeweils dem Indize entnommen werden.

Die zu maximierenden Klassifikationsgenauigkeiten der einzelnen Klassen befinden sich auf der Diagonalen von links oben nach rechts unten (siehe Kapitel 2.5). Die zu minimierenden Fehlerraten sind auf der Nebendiagonalen aufgetragen. Die Berechnung der Klassifikationsgenauigkeiten und der Fehlerraten basiert auf den einzelnen Pixelergebnissen (siehe Kapitel 3.3.3).

Spaltenweise gibt der erste Wert jeder Zelle das Minimum an und der zweite Wert den 25% Punkt. In der Mitte der Zelle ist der Medianwert, welcher auch als 50% Punkt interpretiert werden kann, fett dargestellt. Danach folgt der 75% Punkt und das Maximum der Klassifikation. Diese Werte basieren auf allen fünf Wiederholungen, wobei jeweils 350 Bildausschnitte mit sichtbarem Krankheitsbefall der Größe 64×64 klassifiziert wurden. Die Bildausschnitte stammen dabei wiederum von 55 *Cercospora*-Datensätzen und 73 *Uromyces*-Datensätzen ab. Aufgrund der geringen Anzahl an *Uromyces*-Pixeln in einem Bildausschnitt wurden zur Berechnung der Klassifikationsraten jeweils die Klassifikationsergebnisse von zehn Bildausschnitten zusammengefasst.

Wie in Tabelle 4.2 zu sehen ist, wurden in dem Experiment ohne Berücksichtigung der Nachbarschaftsinformation für die Klasse Gesund im Median eine Klassifikationsgenauigkeit von 84% erreicht, für die Blattfleckenkrankheit *Cercospora beticola* eine Genauigkeit von 74% und für den Braunrost *Uromyces betae* sogar nur eine Genauigkeit in Höhe von 59%. Die Hinzunahme der Nachbarschaftsinformation erhöhte diese Klassifikationsgenauigkeiten ganz erheblich. So konnte für die Klasse Gesund im Median eine Klassifikationsgenauigkeit von 91%, für die Klasse *Cercospora beticola* 81% und für die Klasse *Uromyces betae* 76% erreicht werden.

4.2.3 Optimierung der Bayesklassifikation

Das Ziel der folgenden Experimente ist es, den Bayesklassifikator in Kapitel 2.3 für die Klassifikation von Blattkrankheiten anzupassen. Zu diesem Zweck muss eine geeignete Likelihood- und Gewichtungsfunktion sowie eine möglichst optimale Festlegung der *a-priori* Wahrscheinlichkeit gefunden werden. Zu jedem dieser drei Punkte wird im Folgenden ein Experiment vorgestellt. Die Experimente bauen dabei aufeinander

<i>k</i> -Nächste-Nachbarn Klassifikation						
%	$\hat{h}$		$\hat{i}_{\text{cerc}}$		$\hat{i}_{\text{urom}}$	
$\tilde{h}$ $S \approx 6770000$ px	47.42	68.45	0.24	0.15	2.01	0.52
	76.84	87.49	2.07	2.42	8.68	3.10
	83.65	91.26	4.09	4.59	12.54	4.29
	88.54	93.83	6.29	7.24	18.37	6.22
	97.60	98.86	12.74	19.07	44.71	13.63
$\tilde{i}_{\text{cerc}}$ $S \approx 379500$ px	0.41	0.92	43.50	56.69	1.98	1.36
	2.98	4.52	68.07	73.79	14.06	7.64
	6.39	8.47	74.05	80.95	18.23	9.79
	11.08	14.35	81.70	86.82	21.56	12.71
	40.96	31.54	96.99	96.13	34.87	22.80
$\tilde{i}_{\text{urom}}$ $S \approx 18500$ px	3.41	1.29	1.64	0.92	30.53	36.96
	11.40	5.90	12.27	5.83	53.67	60.67
	16.31	9.21	20.47	14.01	59.25	75.85
	24.53	18.53	27.52	23.11	70.25	84.55
	50.00	46.65	52.84	43.65	81.72	94.51

Tabelle 4.2: Konfusionsmatrix für die k NN-Klassifikation mit und ohne Nachbarschaftsinformation in den Merkmalen und $k = 5$. In allen Tabellen ist die Klasse 'Gesund' als h bezeichnet und die 'kranken' Klassen mit i . Die jeweilige Krankheit kann anhand des Indizes identifiziert werden. Ein geschätzter Wert wird angegeben durch $\hat{h}$ und ein Ground Truth Wert durch $\tilde{h}$. Der erste Wert in jeder Zeile und Zelle gibt das Klassifikationsergebnis ohne die Nachbarschaftsinformation und der zweite Wert das Ergebnis mit Nachbarschaftsinformation an. Spaltenweise ist in jeder Zelle zuerst das Minimum angegeben, gefolgt von dem 25%-Punkt. In der Mitte steht fett gedruckt der Medianwert und zuletzt kommen der 75%-Punkt und das Maximum. In der linken Spalte unter den Ground Truth Labels ist die ungefähre Anzahl an klassifizierten Pixel angegeben auf denen das Ergebnis beruht (siehe Bauer *et al.*, 2011).

auf, d.h. die Erkenntnisse eines vorhergehenden Experimentes werden in dem darauf folgenden direkt mitverwertet.

So basieren beispielsweise die folgenden Experimente alle auf dem erweiterten Merkmalsvektor, da im vorigen Experiment 4.2.2 festgestellt wurde, dass die Klassifikationsgenauigkeiten bei einer Erweiterung des Merkmalsvektors um die Nachbarschaftsinformation deutlich höher liegen.

Das erste Experiment in Kapitel 4.2.4 dient zur Optimierung der Likelihood-Funktion. Für die Verteilung der Daten in der Likelihood-Funktion wird ein Gaußsches Mischmodell angenommen. Ziel dieses Experimentes ist es, die Anzahl der Verteilungen pro Klasse individuell abzuschätzen. Nach der Ermittlung der optimalen Anzahl an Verteilungen erfolgt dann im zweiten Schritt ein Experiment zur Bestimmung der *a-priori* Wahrscheinlichkeiten (siehe Kapitel 4.2.5). Das dritte Experiment dient schließlich zur Anpassung der Gewichtungsfunktion.

Die hier vorgestellten Experimente basieren alle auf denselben Test- und Referenzdaten wie das vorige Experiment 4.2.1. Desweiteren setzen alle drei Experimente jeweils 70 der 700 Bildausschnitt zum Training ein und klassifizieren anschließend die restlichen 630 Bildausschnitte. Jedes Experiment wurde völlig unabhängig voneinander zehnmal wiederholt.

4.2.4 Bestimmung der Anzahl der Verteilungen pro Klasse

Das Ziel dieses Experimentes war es, festzustellen, welche Anzahl an Verteilungen bzw. Clustern c der angenommenen Gaußschen Mischverteilung pro Klasse die höchsten Klassifikationsgenauigkeiten liefert. Dazu wurde das Experiment für $c = 1, \dots, 5$ durchgeführt, wobei für jede Klasse l die gleiche Anzahl an Verteilungen c angenommen wurde. Für die Schätzung der Parameter des Gaußschen Mischmodells wurde der Expectation-Maximization- (EM) Algorithmus (siehe Kapitel 2.3) eingesetzt. Die Bestimmung der Näherungswerte für die Parameter α_c , μ_c und Σ_c erfolgte zufällig. Um eine Beeinflussung der Ergebnisse durch *a-priori* Wahrscheinlichkeit oder die Gewichtungsfunktion auszuschließen, wurde in diesem Experiment eine ML-Klassifikation (siehe Kapitel 2.3.3) durchgeführt.

In Abbildung 4.4 ist die Entwicklung der Klassifikationsgenauigkeiten in Abhängigkeit von der verwendeten Anzahl an Verteilungen dargestellt. Auf der x -Achse ist hier die Anzahl an Verteilungen aufgetragen und auf der y -Achse die Klassifikationsgenauigkeit in %. Die in dieser Abbildung angegebenen Klassifikationsgenauigkeiten sind die Medianwerte, welche jeweils bei der Klassifikation der 6300 Bildausschnitte erzielt wurden. Die Bildausschnitte basieren wiederum auf 55 *Cercospora*-Datensätzen und 73 *Uromyces*-Datensätzen. Wie man sehen kann, nimmt die Klassifikationsgenauigkeit

- des Braunrostes i_{urom} mit zunehmender Anzahl an Verteilungen stetig ab.
- des gesunden Blattes h von einem zu zwei Verteilungen von 95.07% auf 96.68% zu und bleibt danach relativ konstant.

- der Blattfleckenkrankheit i_{cerc} zu, bis die Verteilung durch drei Cluster angenähert wird und bleibt dann ebenfalls relativ konstant.

Die Abbildungen 4.5 und 4.6 zeigen die Entwicklung des Medians der Trainings- und Testzeiten eines $64 \times 64 \times 20$ großen Bildausschnittes bei einer Clusteranzahl von eins bis fünf pro Klasse. Auf der x -Achse ist wiederum die Clusteranzahl aufgetragen und auf der y -Achse in diesem Fall jeweils die benötigte Zeit in Sekunden. Wie man aus diesen Abbildungen entnehmen kann steigt die benötigte Zeit zum Training von knapp 50 Sekunden bei einem Cluster auf fast 1300 Sekunden bei fünf Clustern. Ebenso nimmt die Testzeit eines Bildausschnittes von 0.02 Sekunden bei einem Cluster auf gut 0.09 Sekunden bei fünf Clustern erheblich zu. Aus diesem Grunde sollte die Anzahl an Verteilungen pro Klasse so gering wie möglich gehalten werden, solange nicht eine effizientere Umsetzung der Algorithmen zur Verfügung steht. Es sollte daher für jede Klasse diejenige Anzahl an Verteilungen gewählt werden, für die die erzielbare Klassifikationsgenauigkeit möglichst hoch und die dazu notwendige Clusteranzahl möglichst gering ist.

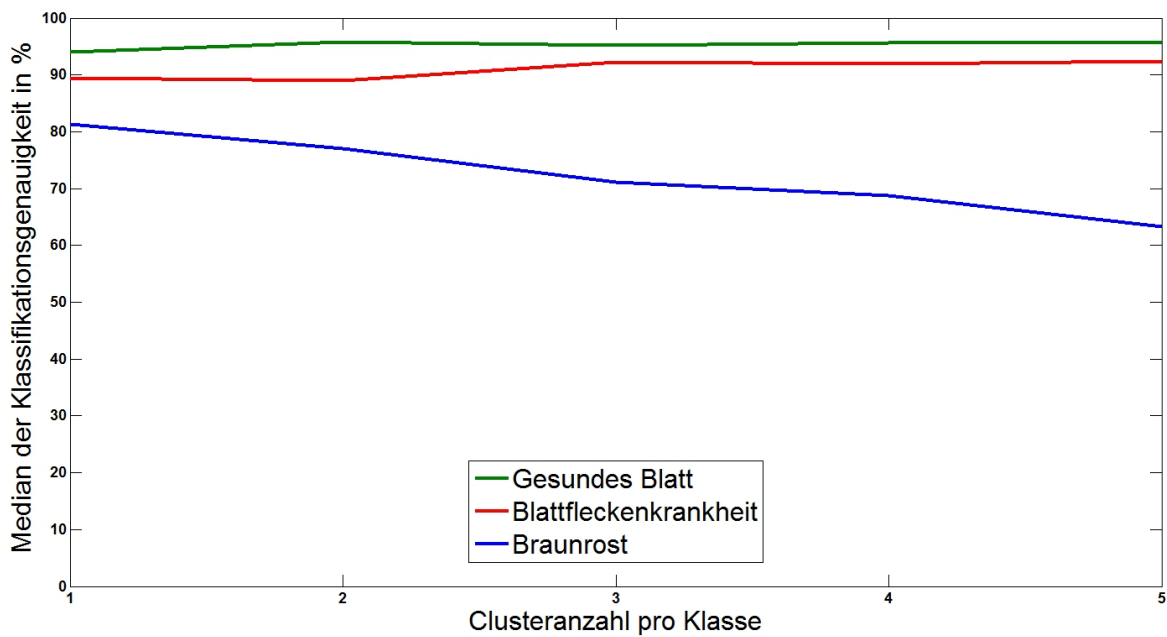


Abbildung 4.4: Median der Klassifikationsgenauigkeit in % in Abhängigkeit von der Clusteranzahl pro Klasse. Auf der x -Achse ist die Clusteranzahl pro Klasse aufgetragen. In diesem Experiment hatten alle drei Klassen immer die gleiche Clusteranzahl. Die y -Achse gibt den Medianwert der Klassifikationsgenauigkeiten in % an. Grün steht für gesunde Blattfläche, rot für die Blattfleckenkrankheit und blau für den Braunrost.

Folgende Anzahl an Verteilungen bieten sich dann für die drei Klassen an:

- Gesunde Blattfläche: 2

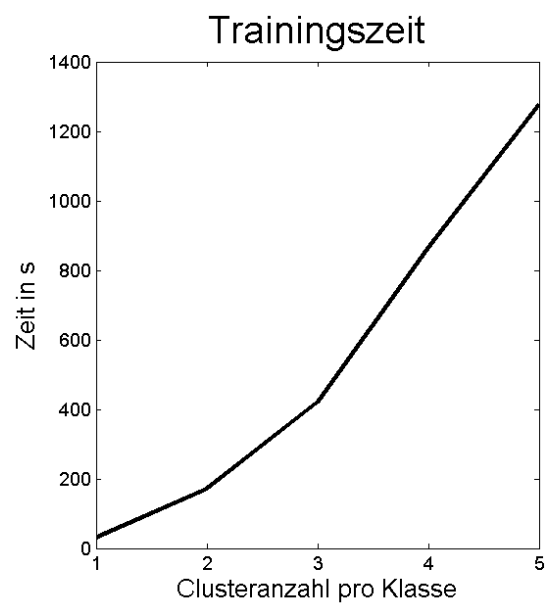


Abbildung 4.5: Median der Trainingszeiten in Sekunden in Abhängigkeit von der Clusteranzahl pro Klasse.

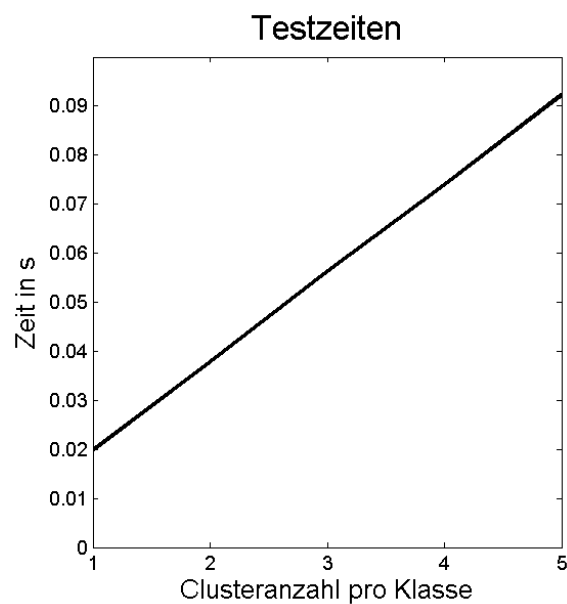


Abbildung 4.6: Median der Testzeiten in Sekunden in Abhängigkeit von der Clusteranzahl pro Klasse.

- *Cercospora beticola*: 3
- *Uromyces betae*: 1

4.2.5 Bestimmung der *a-priori* Wahrscheinlichkeiten

Nachdem im letzten Experiment die beste Anzahl der Verteilungen für jede Klasse bestimmt wurde, soll nun untersucht werden, inwieweit die *a-priori* Wahrscheinlichkeit das Klassifikationsergebnis beeinflusst. Zu diesem Zweck wird eine Maximum *a-posteriori* Klassifikation durchgeführt, wobei die Likelihoodfunktion anhand der im vorigen Experiment gemachten Erkenntnisse aufgestellt wird.

Die *a-priori* Wahrscheinlichkeiten der Klasse h und i_{cerc} bzw. i_{urom} hängen direkt voneinander ab, denn wenn die befallene Blattfläche steigt, verringert sich gleichzeitig die gesunde Blattfläche. Wieviel Blattfläche jedoch an einem einzelnen Blatt tatsächlich noch gesund ist bzw. von einer der beiden Krankheiten befallen ist, ist vorab nur schwer feststellbar. Selbst wenn man mit Hilfe von einigen Trainingsblättern die dortige *a-priori* Wahrscheinlichkeit bestimmt, lässt sich diese nicht unbedingt auf die anderen Blätter übertragen, da der Ausbruchzeitpunkt und die Entwicklungsgeschwindigkeit der Krankheit auf jedem Blatt individuell sehr unterschiedlich sein kann.

Eine Idee wäre es daher, die *a-priori* Wahrscheinlichkeit für jedes zu klassifizierende Blatt durch eine Vorklassifikation des jeweiligen Blattes zu schätzen. Die auf diese Weise geschätzte *a-priori* Wahrscheinlichkeit wäre natürlich nicht ganz korrekt. Daher soll in diesem Experiment die Auswirkung der Wahl der *a-priori* Wahrscheinlichkeit auf das Klassifikationsergebnis untersucht werden.

Aufstellung verschiedener *a-priori* Wahrscheinlichkeiten

Für die Analyse der *a-priori* Wahrscheinlichkeiten sind für jeden Bildausschnitt acht Klassifikationen mit jeweils unterschiedlichen *a-priori* Wahrscheinlichkeiten durchgeführt worden. Die acht verschiedenen *a-priori* Wahrscheinlichkeiten wurden folgendermaßen aufgestellt:

- Zunächst wurde für jede Klasse dieselbe *a-priori* Wahrscheinlichkeit angenommen, so dass im Prinzip eine ML-Klassifikation durchgeführt wurde. Dieses Klassifikationsergebnis diente dann zur Bestimmung der *a-posteriori* Wahrscheinlichkeit von jeder Klasse auf dem jeweiligen Bildausschnitt. Diese *a-posteriori* Wahrscheinlichkeit wurde dann im nächsten Schritt als *a-priori* Wahrscheinlichkeit für eine MAP-Klassifikation verwendet.
- In diesem zweiten Experiment wurde die zuvor in der ML-Klassifikation geschätzte *a-posteriori* Wahrscheinlichkeit als *a-priori* Wahrscheinlichkeit in der MAP-Klassifikation verwendet.

4 Experimente

- Für das dritte Experiment wurden die *a-priori* Wahrscheinlichkeiten anhand des Referenzbildes bestimmt und repräsentierten damit die wirklich korrekte Verteilung der Klassen auf dem jeweiligen Bildausschnitt.
- In den folgenden fünf Experimenten wurde diese korrekte *a-priori* Wahrscheinlichkeit schrittweise verändert. Die *a-priori* Wahrscheinlichkeit der Klasse Gesund, h , wurde in 2er Potenzschritten reduziert und die Wahrscheinlichkeit der auf dem jeweiligen Bildausschnitt vertretenen Blattkrankheit i im selben Maße erhöht:

$$P_{\nu+1}(h) = P_{\nu}(h) - \frac{2^j}{100} \text{ für } j = 1, \dots, 5, \quad (4.1)$$

$$P_{\nu+1}(i) = P_{\nu}(i) + \frac{2^j}{100} \text{ für } j = 1, \dots, 5. \quad (4.2)$$

Da die *a-priori* Wahrscheinlichkeit $P(h)$ im Median 98.93 % beträgt und minimal 41.99 % erreicht, ist dieses Vorgehen problemlos möglich, ohne dass $P_{\nu+1}(h)$ negativ wird.

Klassifikationsgenauigkeiten bei der Wahl verschiedener *a-priori* W.

In Abbildung 4.7 ist die Entwicklung der Klassifikationsgenauigkeiten für die acht verschiedenen *a-priori* Wahrscheinlichkeiten im Median dargestellt. Die Klassifikationsgenauigkeiten basieren dabei, wie im vorigen Experiment, auf 6300 Bildausschnitten aus 55 *Cercospora*- und 73 *Uromyces*-Datensätzen. Von links nach rechts ist in Abbildung 4.7 auf der x -Achse zuerst das Ergebnis der ML-Klassifikation aufgetragen. Danach folgt das Ergebnis der MAP-Klassifikation, wo als *a-priori* Wahrscheinlichkeit das *a-posteriori* Ergebnis der ML-Klassifikation verwendet wurde. Als nächstes folgt das Ergebnis der MAP-Klassifikation mit der korrekten *a-priori* Wahrscheinlichkeit, welche anhand des Labelbildes des jeweiligen Bildausschnittes ermittelt wurde. Danach folgen die Ergebnisse der MAP-Klassifikation, wo die korrekte *a-priori* Wahrscheinlichkeit der Klasse h in 2er Potenzschritten reduziert wurde und im Gegenzug die *a-priori* Wahrscheinlichkeit der auf dem jeweiligen Bildausschnitt vertretenen Blattkrankheit erhöht wurde.

Wie man in dieser Abbildung sehen kann, sind die Klassifikationsgenauigkeiten der drei Klassen bei Verwendung der geschätzten *a-priori* Wahrscheinlichkeit aus der *a-posteriori* Wahrscheinlichkeit der ML-Klassifikation und der korrekten *a-priori* Wahrscheinlichkeit sehr ähnlich. Die genauen Ergebnisse dieser beiden Experimente sind in Tabelle 4.3 dargestellt.

Das beste Klassifikationsergebnis mit Genauigkeiten zwischen 93% und 97% wurden bei einer Reduktion der korrekten *a-priori* Wahrscheinlichkeit für die Klasse h und einer Erhöhung der jeweiligen Blattkrankheit i_{cerc} oder i_{urom} um $\frac{2^5}{100}$ erreicht.

Die Ergebnisse dieser Analyse zeigen, dass die Klassifikationsgenauigkeiten bei Einsatz der geschätzten und der korrekten *a-priori* Wahrscheinlichkeit nur marginal voneinander abweichen. Damit kann also durchaus die *a-priori* Wahrscheinlichkeit für MAP-Klassifikationen vorab durch eine ML-Klassifikation geschätzt werden.

Insgesamt gesehen hat die Wahl der *a-priori* Wahrscheinlichkeit einen sehr großen Einfluss auf das Klassifikationsergebnis, da durch die Veränderung der *a-priori* Wahrscheinlichkeit die Klassen unterschiedlich gewichtet werden. Eine Erhöhung der an sich eher geringen *a-priori* Wahrscheinlichkeiten der Blattkrankheiten verbessern dadurch deren Klassifikationsgenauigkeit erheblich. Maximal konnten so Klassifikationsgenauigkeiten von 97% bei der Klasse h und 93% bei der Klasse i_{cerc} sowie 94% bei der Klasse i_{urom} erreicht werden.

Es hat sich somit gezeigt, dass eine höhere Gewichtung der seltenen Klassen i_{cerc} und i_{urom} deutlich bessere Klassifikationsgenauigkeiten liefern. Daher wollen wir in einer weiteren Untersuchung eine optimale Kostenfunktion suchen und damit eine Bayesklassifikation mit Risikominimierung durchführen.

4.2.6 Aufstellung der Gewichtungsfunktion für die Bayesklassifikation

Im vorigen Experiment 4.2.5 hat sich gezeigt, dass eine höhere Gewichtung der eher seltenen Blattkrankheiten zu einer besseren Klassifikationsgenauigkeit dieser Klassen führt. Das folgende Experiment dient daher zur Analyse, welche Kostenmatrix bei einer Bayesklassifikation mit Risikominimierung von Blattkrankheiten die besten Ergebnisse liefert.

Aufstellung verschiedener Gewichtungsfunktionen

Im vorigen Experiment konnten ohne Berücksichtigung von Ground Truth Informationen die besten Klassifikationsgenauigkeiten bei einer ML-Klassifikation erreicht werden. Aus diesem Grund wird im Folgenden die Kostenmatrix mit Gewichten initialisiert, welche anhand der *a-priori* Wahrscheinlichkeiten der einzelnen Klassen bestimmt wurden, um damit zunächst eine ML-Klassifikation zu simulieren. Diese Initialisierung der Kostenmatrix wird hier Basiskostenmatrix c_{basis} genannt.

Für die Erstellung der Basiskostenmatrix c_{basis} sind zwei verschiedene *a-priori* Wahrscheinlichkeiten $P(l)$ verwendet worden:

1. Die *a-priori* Wahrscheinlichkeit $P(l)$ ist anhand der Ground Truth Daten berechnet worden. Die Basiskosten, welche auf dieser korrekten *a-priori* Wahrscheinlichkeit beruhen, bezeichnen wir im Folgenden als prior-Basiskostenmatrix.
2. Die Trainingsdaten sind mit einer ML-Klassifikation klassifiziert und die daraus resultierenden *a-posteriori* Wahrscheinlichkeiten sind zur Berechnung der Basis-kostenmatrix verwendet worden. Die daraus resultierende Basiskostenmatrix wird dementsprechend im Folgenden als posterior-Basiskostenmatrix bezeichnet.

4 Experimente

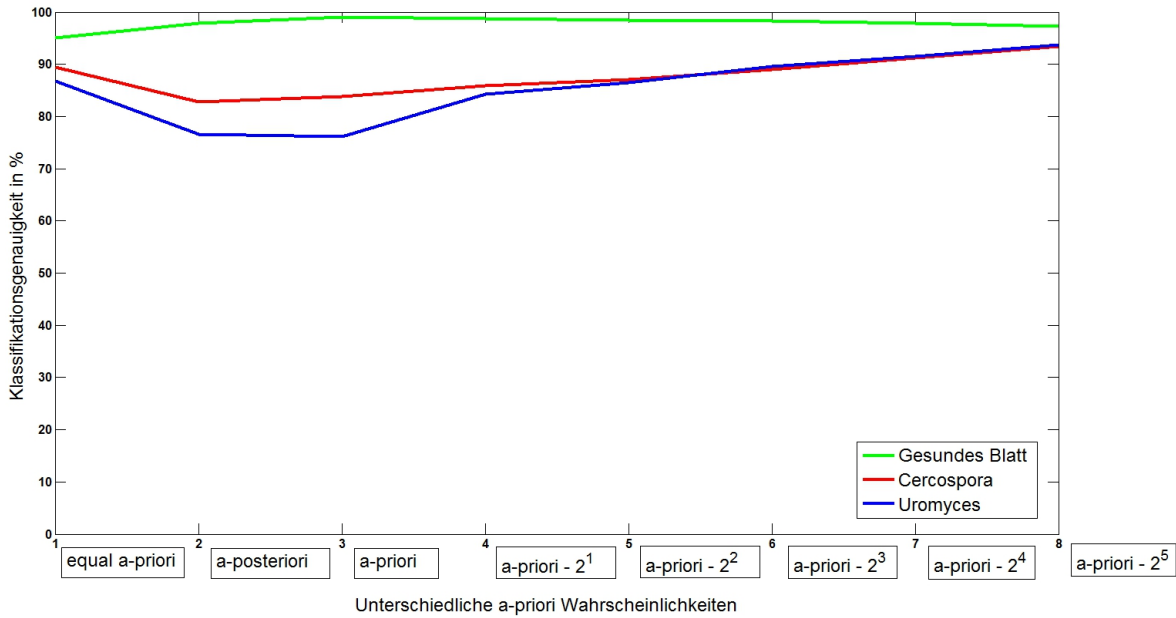


Abbildung 4.7: Median der Klassifikationsgenauigkeit in % in Abhängigkeit von der gewählten *a-priori* Wahrscheinlichkeit pro Klasse. Auf der *x*-Achse sind die verschiedenen *a-priori* Wahrscheinlichkeiten aufgetragen: Als erstes ist der Fall aufgetragen, dass alle Klassen die gleiche Auftretswahrscheinlichkeit haben. An zweiter Stelle folgt die *a-posteriori* Wahrscheinlichkeit, welche mittels einer ML-Klassifikation der Testdaten ermittelt wurde. An dritter Stelle kommt die korrekte *a-priori* Wahrscheinlichkeit, welche anhand der Ground Truth Daten ermittelt wurde. Diese korrekte *a-priori* Wahrscheinlichkeit wurde nun verändert, indem die Wahrscheinlichkeit der Klasse h in 2er Potenzen erhöht wurde und die Wahrscheinlichkeit der kranken Blattfläche entsprechend verringert wurde. Da sich die Wahrscheinlichkeiten der *a-priori* Wahrscheinlichkeit zwischen 0 und 1 bewegen, wurden die 2er Potenzen ebenfalls durch 100 dividiert. Aus Platzmangel ist dies allerdings in der Grafik bei den Abszissenwerten nicht aufgeführt. Auf der *y*-Achse ist der Medianwert der Klassifikationsgenauigkeit in % aufgetragen.

MAP-Klassifikation						
%	$\hat{h}$		$\hat{i}_{\text{cerc}}$		$\hat{i}_{\text{urom}}$	
$\tilde{h}$	70.74	73.58	0.00	0.00	0.02	0.00
	96.26	97.56	0.13	0.00	0.44	0.00
	97.90	98.97	0.83	0.44	0.87	0.04
	99.05	99.52	2.46	2.44	1.47	0.49
	99.96	99.97	28.87	26.42	5.97	2.07
$\tilde{i}_{\text{cerc}}$	0.01	0.05	17.89	24.88	0.00	0.00
	7.09	8.84	72.38	74.64	0.95	0.00
	13.85	16.17	82.83	83.83	2.33	0.00
	23.34	25.36	90.30	91.16	4.74	0.00
	73.23	75.12	99.90	99.95	20.80	0.00
$\tilde{i}_{\text{urom}}$	0.00	1.00	0.00	0.00	6.43	10.04
	7.16	12.32	0.27	0.00	56.46	57.13
	15.15	23.81	2.72	0.00	76.60	76.19
	27.91	42.87	9.62	0.00	89.63	87.68
	83.88	89.96	65.21	0.00	100.00	99.00

Tabelle 4.3: Konfusionsmatrix für die Maximum *a-posteriori* Klassifikation unter Verwendung einer vorab durch eine ML-Klassifikation geschätzte *a-priori* Wahrscheinlichkeit (Experiment 2) und die korrekte *a-priori* Wahrscheinlichkeit (Experiment 3). Der Aufbau der Tabelle ist prinzipiell der gleiche wie in Tabelle 4.2. Nur dass hier der erste Wert in jeder Zeile und Zelle das Klassifikationsergebnis des Experimentes 2 zeigt und der zweite Wert das Ergebnis des Experimentes 3.

4 Experimente

Die so erstellten Basiskostenmatrizen wurden dann in einem zweiten Schritt anhand des in Kapitel 3.3.4 vorgestellten Verfahren zu der Gewichtungsfunktion w optimiert. Die Gewichtungsfunktion w , welche auf der prior-Basiskostenmatrix basiert, bezeichnen wir im Folgenden als prior-Kostenmatrix und die andere entsprechend als posterior-Kostenmatrix.

Es wurden damit vier Experimente mit folgenden Gewichtungsfunktionen durchgeführt:

- prior-Basiskostenmatrix, welche aufgrund der korrekten *a-priori* W. erstellt worden ist.
- posterior-Basiskostenmatrix, welche auf der in der ML-Klassifikation der Trainingsdaten bestimmten *a-posteriori* W. beruht.
- prior-Kostenmatrix: optimierte prior-Basiskostenmatrix.
- posteriori-Kostenmatrix: optimierte posterior-Basiskostenmatrix.

Die Klassifikation erfolgte mit dem in Kapitel 3.3.4 vorgestellten Bayesklassifikator. Aufgrund des Ergebnisses des Experimentes zur Bestimmung der *a-priori* W. (siehe 4.2.5) ist die *a-priori* Wahrscheinlichkeit für jeden Test-Bildausschnitt individuell durch Vorklassifikation mit einer ML-Klassifikation geschätzt worden. Für die Likelihood-Funktion sind analog zu den Ergebnissen des Experimentes in Kapitel 4.2.4 Gaußsche Mischmodelle mit zwei Verteilungen für die Klasse Gesund, drei Verteilungen für die Blattfleckenkrankheit und einer Verteilung für den Braunrost erstellt worden.

Klassifikationsgenauigkeiten bei der Wahl verschiedener Gewichtungsfunktionen

Klassifikationsgenauigkeiten bei verschiedenen Gewichtungsfunktionen			
%	$\hat{h}$	$\hat{i}_{\text{cerc}}$	$\hat{i}_{\text{urom}}$
prior-Basiskostenmatrix	90.79	88.11	93.93
prior-Kostenmatrix	88.92	86.39	94.73
posterior-Basiskostenmatrix	95.05	91.89	84.70
posterior-Kostenmatrix	94.44	91.26	85.97

Tabelle 4.4: Mediane der erzielten Klassifikationsgenauigkeiten für jede Klasse bei Wahl der verschiedenen Gewichtungsfunktionen.

4.2 Entwicklung der pixelweisen Klassifikationsstrategie

Eine Übersicht über die Klassifikationsgenauigkeiten der einzelnen Klassen bei der Wahl der vier verschiedenen Gewichtungsfunktionen ist in Tabelle 4.4 dargestellt. Die ausführlichen Ergebnisse der letzten beiden Experimente sind zusätzlich in Tabelle 4.5 aufgeführt. Wie schon in den beiden vorigen Experimenten, beruhen auch die Ergebnisse in diesem Experiment auf der Klassifikation von insgesamt jeweils 6300 Bildausschnitten aus 55 *Cercospora*- und 73 *Uromyces*-Datensätzen.

Wie den Pixelzahlen S in Tabelle 4.5 entnommen werden kann, gehören die meisten Pixel zur Klasse gesund (94.45%), ein kleiner Teil zur Klasse *Cercospora* (5.30%) und fast keine zur Klasse *Uromyces* (0.25%). Aus diesem Grund sollte eine derart hohe Fehlklassifikation der Klasse gesund, wie sie bei Verwendung der prior-Basiskostenmatrix und der prior-Kostenmatrix auftritt (siehe Tabelle 4.4), vermieden werden. Daher ist die Berechnung der Gewichtungsfunktion anhand der posterior-Verteilung zu bevorzugen.

Durch die Optimierung der posterior-Basiskostenmatrix kann die Klassifikationsgenauigkeit von *Uromyces* noch um ein gutes Prozent erhöht werden, die Genauigkeit der anderen beiden Klassen sinkt dafür im Gegenzug um ein gutes halbes Prozent. Da aber auch die seltene Blattkrankheit *Uromyces betae* erkannt werden sollte, scheint der letzte Fall der beste Kompromiss zu sein.

Adaptive Bayesklassifikation						
%	$\hat{h}$		$\hat{i}_{\text{cerc}}$		$\hat{i}_{\text{urom}}$	
$\tilde{g}$ $S \approx 24400000$ px	62.83	60.96	0.00	0.00	0.01	0.03
	91.21	90.44	0.29	0.27	1.14	1.41
	95.05	94.44	2.02	2.00	2.22	2.76
	97.73	97.33	5.34	5.25	3.59	4.38
	99.85	99.82	32.15	32.98	16.55	21.61
$\tilde{i}_{\text{cerc}}$ $S \approx 1370000$ px	0.00	0.00	53.02	54.22	0.00	0.00
	1.17	1.03	87.23	86.76	1.43	1.76
	3.79	3.68	91.89	91.26	2.91	3.75
	8.25	7.66	96.81	96.47	5.49	6.52
	41.84	39.62	99.83	99.86	20.87	21.84
$\tilde{i}_{\text{urom}}$ $S \approx 64600$ px	0.00	0.00	0.00	0.00	26.24	27.36
	2.03	1.72	0.65	0.46	67.07	69.86
	4.77	3.94	4.06	3.41	84.70	85.97
	14.29	12.70	13.38	12.75	94.20	95.42
	72.64	72.64	66.36	66.36	100.00	100.00

Tabelle 4.5: Der Aufbau dieser Tabelle ist prinzipiell identisch zu den vorigen Konfusionsmatrizen (siehe Tabelle 4.2). Hier zeigt jedoch der erste Wert in jeder Zeile und Zelle jeweils das Ergebnis für das vorletzte Experiment, basierend auf der posterior-Basiskostenmatrix. Der zweite Wert repräsentiert das Ergebnis des letzten Experiments (Bauer *et al.*, 2011), basierend auf der optimierten posterior-Kostenmatrix.

4.3 Glättung der pixelweisen Klassifikation

Nach der Entwicklung der pixelweisen Klassifikationsstrategie folgt nun die Glättung des pixelweisen Ergebnisses mit einem Majoritätsfilter oder mit bedingten Markoffschen Zufallsfeldern. Im nächsten und letzten Schritt erfolgt dann eine regionenbasierte Klassifikation. Die regionenbasierte Klassifikation kann dabei nur auf ganzen Blättern angewandt werden, da ansonsten Fehler an den Rändern der Bildausschnitte durch die Beschneidung von Blattflecken auftreten. Aus diesem Grunde wird auch die Glättung bereits auf ganzen Blättern durchgeführt, im Gegensatz zu den in den vorigen Experimenten eingesetzten Bildausschnitten.

Die Durchführung des Trainings und der anschließenden Klassifikation basierte auf den im Kapitel 3.3.2 vorgestellten Orthofotos mit dem erweiterten Merkmalsvektor unter Einsatz des im vorigen Experiments ermittelten und in Kapitel 3.3.4 beschriebenen adaptiven Bayesklassifikators. Das Training des Bayesklassifikators erfolgte in diesem Falle so, dass von jedem Aufnahmedatum 0.2 der *Cercospora*-Bilder und 0.3 der *Uromyces*-Bilder als Trainingsbilder verwandt wurden. Aus den Trainingsbildern wurden pro Klasse mindestens 500 und höchstens 15000 Pixel zufällig ausgewählt. D.h., wenn für eine Klasse zuwenig Pixel in den Trainingsbildern enthalten waren, wurden die Trainingsbilder erneut zufällig ausgewählt, bis für alle Klassen die Mindestpixelzahl erreicht worden war. Für den Fall, dass für eine Klasse mehr als 15000 Pixel vorhanden waren, wurde aus dieser größeren Menge 15000 Pixel zufällig ausgewählt.

Die restlichen 180 Bilder wurden im Anschluß an das Training klassifiziert. Als Ergebnis der Bayesklassifikation lagen sogenannte Labelbilder vor, in welchem jedem Pixel einer Zahlenwert zwischen 0 und 3 zugewiesen wurde. Die Zahlenwerte haben folgende Bedeutung

- 0: Hintergrund
- 1: Gesunde Blattfläche
- 2: *Cercospora*
- 3: *Uromyces*

Nachstehend erfolgt die Erläuterung der Experimente zur Glättung dieser Labelbilder.

4.3.1 Anwendung eines Majoritätsfilters oder Markoffscher Zufallsfelder

Zur Glättung wurden zwei Experimente durchgeführt. Im ersten Experiment wurden die aus der Bayesklassifikation resultierenden Labelbilder mit einem Majoritätsfilter geglättet und im zweiten Experiment mit bedingten Markoffschen Zufallsfeldern. Wie in Kapitel 3.3.5 erläutert, berücksichtigt der Majoritätsfilter die Klassen des jeweiligen Pixels selbst und der 4er-Nachbarschaft, d.h. des linken, rechten, oberen und unteren Nachbarn. Er weist dem Pixel diejenige Klasse zu, welche am häufigsten vertreten ist.

Auf diese Weise wurden alle 180 aus der Bayesklassifikation resultierenden Labelbilder geglättet.

Die Schätzung der bedingten Markoffschen Zufallsfelder ist sehr rechenintensiv. Aus diesem Grund wurden die Bilder für die Glättung mit den Zufallsfeldern in sich überlappende Ausschnitte der Größe 96×98 Pixel aufgeteilt. Die Überlappung betrug an allen Seiten jeweils eine Pixelbreite. Bei der vorliegenden Auflösung von 1600×1250 Pixel der Orthofotos resultiert dies in 221 Bildausschnitte. Die Berechnungszeit für einen Bildausschnitt dieser Größe variierte bei Verwendung des Active-Set Algorithmus' zwischen ungefähr 27 Sekunden und maximal fast 15 Minuten auf einem System mit einem Intel Core i7 860 Prozessor und 8 GB RAM. Die einzelnen Bildausschnitte wurden abschließend wieder zu einem Bild zusammengesetzt. Die Klassifikationszeit eines kompletten Bildes betrug damit ca. 3,5 Stunden. Aufgrund dieser langen Berechnungszeit wurden lediglich 12 der 180 Bilder mit dem CRF geglättet.

4.3.2 Vergleich der beiden Glättungsarten

Die Ergebnisse der Glättung mit dem Majoritätsfilter oder den bedingten Markoffschen Zufallsfeldern ist in den Tabellen 4.6 bis 4.11 zu finden. Die Tabelle 4.6 zeigt die minimalen, maximalen und im Median erreichten Klassifikationsgenauigkeiten für alle zwölf mit den bedingten Markoffschen Zufallsfeldern geglätteten Bildern. Die zwölf Bilder setzen sich folgendermaßen zusammen:

- 6 völlig symptomlose Blätter
- 3 mit *Cercospora beticola* infizierte Blätter
 - 1 Blatt mit ganz leichtem Befall
 - 1 Blatt mit mittlerem Befall
 - 1 Blatt mit starkem Befall
- 3 mit *Uromyces betae* infizierte Blätter
 - 1 Blatt mit ganz leichtem Befall
 - 1 Blatt mit mittlerem Befall
 - 1 Blatt mit starkem Befall

Zum besseren Vergleich basieren die Ergebnisse in den drei Tabellen 4.6, 4.7 und 4.8 auf denselben zwölf Bildern. Die Tabelle 4.7 zeigt dabei die Klassifikationsgenauigkeiten für eine Glättung mit dem Majoritätsfilter und die Tabelle 4.8 enthält zur Referenz die Genauigkeiten für die pixelweise, adaptive Bayesklassifikation, auf welcher die beiden Glättungsverfahren basieren.

Die Tabellen 4.9, 4.10 und 4.11 stellen die Ergebnisse der pixelweisen, adaptiven Bayesklassifikation und der beiden Glättungsverfahren jeweils für genau ein konkretes

Bedingte Markoffsche Zufallsfelder			
%	$\hat{h}$	$\hat{i}_{\text{cerc}}$	$\hat{i}_{\text{urom}}$
$\tilde{h}$	84.62	0.91	0.00
	96.40	2.97	0.16
	98.95	14.26	1.13
$\tilde{i}_{\text{cerc}}$	5.73	92.28	0.00
	6.77	93.23	0.29
	7.01	93.98	0.71
$\tilde{i}_{\text{urom}}$	5.80	0.00	60.71
	7.71	9.41	82.88
	17.86	21.43	94.20

Tabelle 4.6: Konfusionsmatrix für die Glättung des pixelweisen Klassifikationsergebnisses (siehe Tabelle 4.8) mit Hilfe der bedingten Markoffschen Zufallsfelder. Der erste Wert gibt jeweils den minimal erzielten Wert an. Der mittlere, fett hervorgehobene Wert den Median und die letzte Zahl ist der maximal erreichte Wert.

Majoritätsfilterung			
%	$\hat{h}$	$\hat{i}_{\text{cerc}}$	$\hat{i}_{\text{urom}}$
$\tilde{h}$	87.44	0.50	0.00
	97.17	1.90	0.34
	99.12	9.46	3.10
$\tilde{i}_{\text{cerc}}$	8.75	80.64	3.89
	9.61	85.21	6.04
	12.23	86.51	7.13
$\tilde{i}_{\text{urom}}$	5.80	0.00	64.29
	7.82	3.85	88.32
	25.00	10.71	94.20

Tabelle 4.7: Konfusionsmatrix für die Majoritätsfilterung basierend auf der pixelweisen Bayesklassifikation (siehe Tabelle 4.8).

Adaptive Bayesklassifikation			
%	$\hat{h}$	$\hat{i}_{\text{cerc}}$	$\hat{i}_{\text{urom}}$
$\tilde{h}$	86.24	0.57	0.01
	96.65	1.98	0.52
	98.87	9.61	4.15
$\tilde{i}_{\text{cerc}}$	7.99	79.81	5.22
	9.25	83.80	8.21
	10.93	85.52	9.26
$\tilde{i}_{\text{urom}}$	5.80	2.90	64.29
	6.24	4.31	89.46
	17.86	17.86	91.30

Tabelle 4.8: Konfusionsmatrix für die pixelweise, adaptive Bayesklassifikation über alle zwölf Blätter.

Genauigkeiten der beiden Glättungsarten				
Blattfleckenkrankheit				
%		$\hat{h}$	$\hat{i}_{\text{cerc}}$	$\hat{i}_{\text{urom}}$
$\tilde{h}$	Bayes	86.24	9.61	4.15
	Filterung	87.44	9.46	3.10
	CRF	84.62	14.26	1.13
$\tilde{i}_{\text{cerc}}$	Bayes	7.99	83.80	8.21
	Filterung	8.75	85.21	6.04
	CRF	5.73	93.98	0.29

Tabelle 4.9: Vergleich der erzielten Klassifikationsgenauigkeiten für die pixelweise Bayesklassifikation (jeweils die 1. Zeile), eine folgende Glättung mit einem Majoritätsfilter (jeweils die 2. Zeile) und der Anwendung von bedingten Markoffschen Zufallsfeldern, hier CRF genannt (jeweils die 3. Zeile). Die Ergebnisse beziehen sich auf ein mit *Cercospora beticola* infiziertes Blatt.

Genauigkeiten der beiden Glättungsarten				
Braunrost				
%		$\hat{h}$	$\hat{i}_{\text{cerc}}$	$\hat{i}_{\text{urom}}$
$\tilde{h}$	Bayes	98.73	0.79	0.48
	Filterung	98.98	0.72	0.30
	CRF	98.84	0.96	0.21
$\tilde{i}_{\text{urom}}$	Bayes	6.24	4.31	89.46
	Filterung	7.82	3.85	88.32
	CRF	7.71	9.41	82.88

Tabelle 4.10: Der Aufbau dieser Tabelle ist der gleiche wie in der vorigen Tabelle 4.9. Die angegebenen Klassifikationsgenauigkeiten basieren hier allerdings auf einem mit dem Braunrost *Uromyces betae* infizierten Blatt.

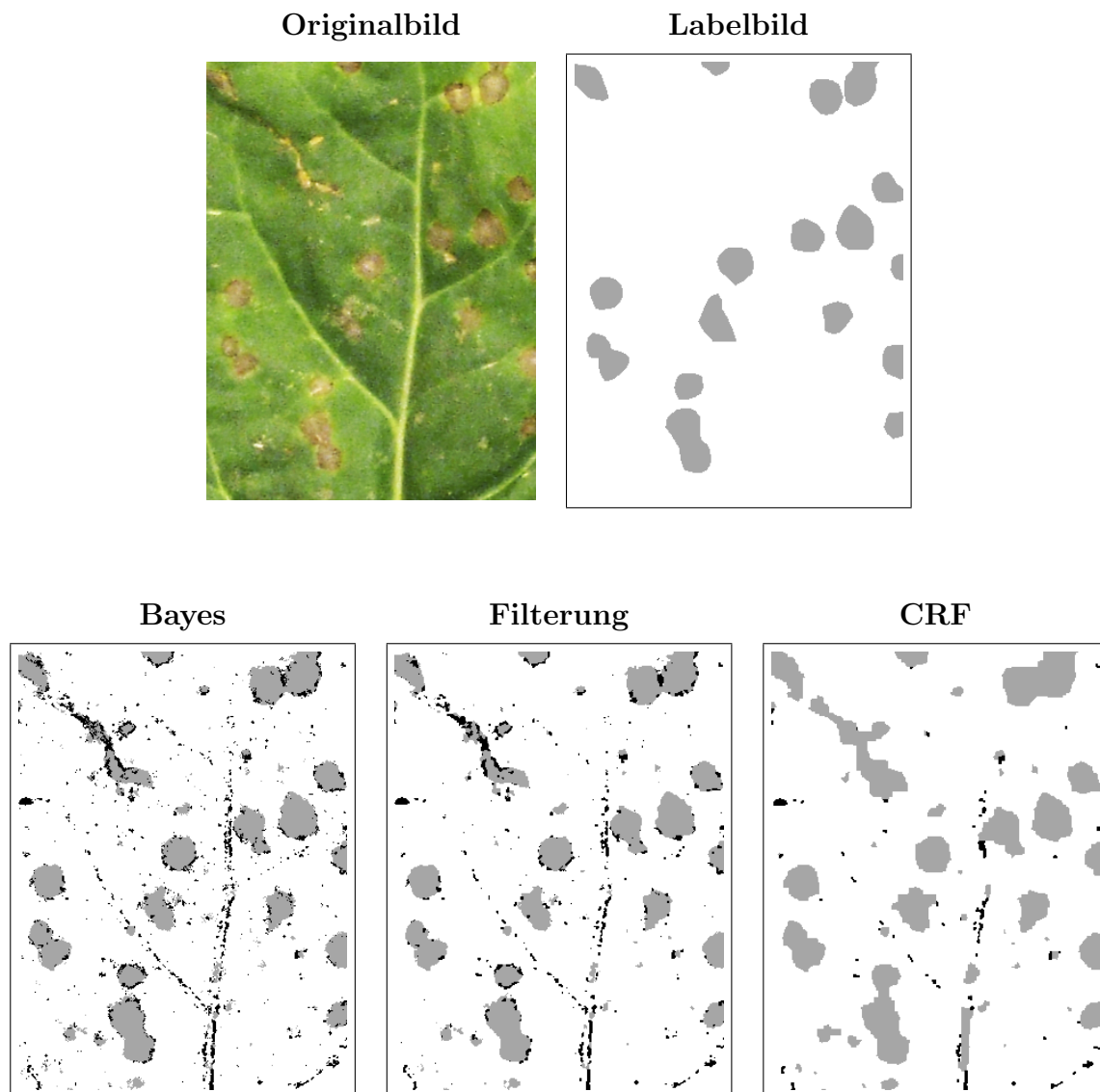


Abbildung 4.8: Klassifikationsvergleich auf einem Blatt mit Befall der Blattfleckenkrankheit. Von links oben nach rechts unten ist zuerst der Originalausschnitt des RGB-Bildes dargestellt, anschließend der dazugehörige Ausschnitt des Ground Truth Labelbildes, gefolgt von dem Klassifikationsergebnis der pixelweisen, adaptiven Bayesklassifikation und der beiden verschiedenen Glättungsvarianten durch den Majoritätsfilter oder die bedingten Markoff'schen Zufallsfelder. Beide Glättungen wurden direkt auf dem Labelbild der pixelweisen Bayesklassifikation ausgeführt. In den Labelbildern repräsentiert Weiß die gesunde Blattfläche, Grau die Blattfleckenkrankheit und Schwarz den Braunrost.

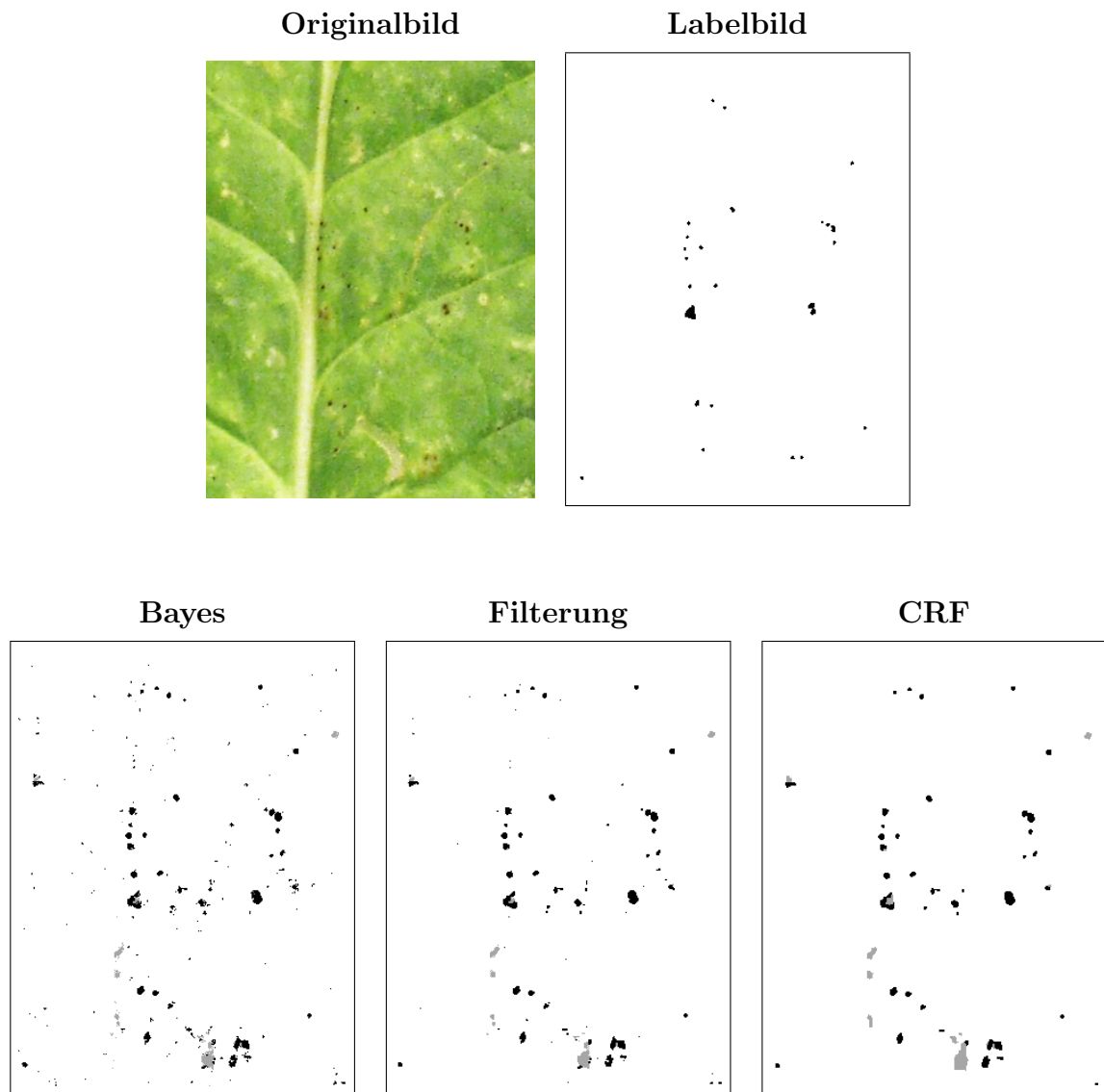


Abbildung 4.9: Klassifikationsvergleich auf einem mit Braunrost befallenen Blatt.

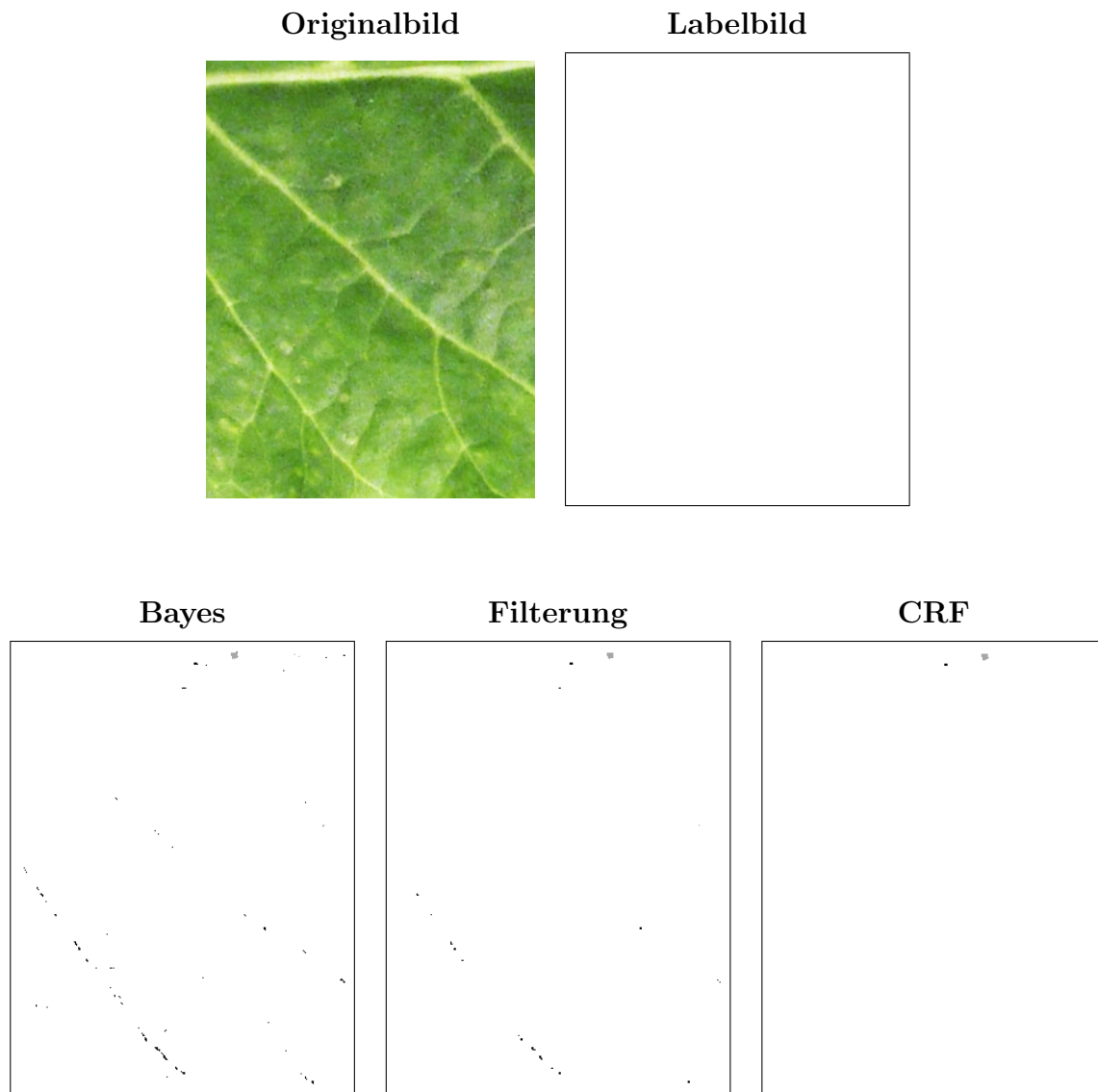


Abbildung 4.10: Klassifikationsvergleich auf einem gesunden Blatt.

Genauigkeiten der beiden Glättungsarten				
Gesund				
%		$\hat{h}$	$\hat{i}_{cerc}$	$\hat{i}_{urom}$
$\tilde{h}$	Bayes	98.62	1.26	0.12
	Filterung	98.75	1.20	0.05
	CRF	98.55	1.44	0.01

Tabelle 4.11: In dieser werden die Ergebnisse der beiden Glättungsarten bezüglich eines gesunden Blattes verglichen.

Blatt gegenüber. Die erste Tabelle 4.9 zeigt die erreichten Klassifikationsgenauigkeiten für ein stark mit *Cercospora beticola* infiziertes Blatt. In der ersten Zeile dieser und der beiden folgenden Tabellen ist jeweils das Ergebnis der pixelweisen, adaptiven Bayesklassifikation angegeben. In der zweiten Zeile folgen die Klassifikationsgenauigkeiten einer Glättung mit dem Majoritätsfilter und in der dritten Zeile steht das Ergebnis der bedingten Markoffschen Zufallsfelder (CRF). Da auf dem Blatt nur die Klassen Gesund h und *Cercospora beticola* i_{cerc} auftreten, sind auch nur diese beiden Klassen auf der linken Seite der Tabelle als Ground Truth Klassen angegeben. Die zweite Tabelle 4.10 gibt die verschiedenen Klassifikationsgenauigkeiten für ein stark mit *Uromyces betae* infiziertes Blatt an und die dritte und letzte Tabelle 4.11 zeigt die Ergebnisse für ein völlig gesundes Blatt.

In den Abbildungen 4.8, 4.9 und 4.10 sind Bildausschnitte von jedem der drei Blätter dargestellt, deren Klassifikationsgenauigkeiten in den Tabellen 4.9, 4.10 und 4.11 angegeben sind. In der ersten Reihe ist jeweils der Ausschnitt des Original-RGB-Bildes und der dazugehörige Ausschnitt des Ground Truth Labelbildes dargestellt. In der Reihe darunter ist von links nach rechts zuerst das Ergebnis der pixelweisen, adaptiven Bayesklassifikation abgebildet, anschließend das Resultat der Majoritätsfilterung und ganz rechts das Resultat der bedingten Markoffschen Zufallsfelder. Die Filterung und die Anwendung der bedingten Markoffschen Zufallsfelder erfolgte jeweils auf dem aus der pixelweisen, adaptiven Bayesklassifikation resultierenden Labelbild. In allen Labelbildern repräsentiert Weiß die gesunde Blattfläche, Grau die Blattfleckenkrankheit *Cercospora beticola* und Schwarz den Braunrost *Uromyces betae*.

Sowohl in den Tabellen 4.6, 4.7 und 4.9 als auch in der dazugehörigen Abbildung 4.8 ist ersichtlich, dass die höchste Klassifikationsgenauigkeit für die Blattfleckenkrankheit *Cercospora beticola* bei Anwendung der bedingten Markoffschen Zufallsfelder erreicht werden kann. So liegt die Klassifikationsgenauigkeit für die Blattfleckenkrankheit auf einem stark befallenen Blatt nach Durchführung der adaptiven Bayesklassifikation bei 84%, bei einer anschließenden Filterung bei 85% und bei Anwendung der bedingten

4 Experimente

Markoffschen Zufallsfelder bei 94% (siehe Tabelle 4.9). Zusätzlich kann dies auch an den minimal und maximal erzielten Werten in den Tabellen 4.6, 4.7 und 4.8 abgelesen werden. In jedem Fall verbessert also die Glättung das Klassifikationsergebnis dieser Klasse.

Durch die sehr geringe Größe der Rostpusteln, welche in der Regel nur wenige Pixel umfassen, führt die Glättung hier bei beiden Verfahren zu einer geringeren Klassifikationsgenauigkeit dieser Klasse. So klassifizierte der adaptive Bayesklassifikator den Braunrost im Median mit einer Genauigkeit von 89%. Diese Genauigkeit reduzierte sich dann bei Anwendung der Majoritätsfilterung auf 88% und bei Einsatz der bedingten Markoffschen Zufallsfelder sogar auf 83%. Die maximal erreichte Klassifikationsgenauigkeit des Braunrostes steigt dagegen bei beiden Verfahren von 91% auf 94% (siehe Tabellen 4.6, 4.7 und 4.8). Die minimal erzielte Klassifikationsgenauigkeit bleibt bei der Majoritätsfilterung mit 64% identisch zum Ergebnis der Bayesklassifikation. Bei Anwendung der bedingten Markoffschen Zufallsfelder sinkt sie um 3% auf 61% ab.

Die Klassifikationsgenauigkeit der Klasse Gesund kann dagegen durch eine Glättung mit dem Majoritätsfilter in allen Fällen erhöht werden. So steigt sie auf einem völlig gesunden Blatt von 98.62% auf 98.75% (siehe Tabelle 4.11), auf einem mit *Cercospora beticola* infizierten Blatt von 86.24% auf 87.44% (siehe Tabelle 4.9) und auf einem mit *Uromyces betae* infizierten Blatt von 98.73% auf 98.98% (siehe Tabelle 4.10). Im Median bedeutet dies eine Steigerung von 96.65% auf 97.17% (siehe Tabellen 4.7 und 4.8).

Bei einer Glättung durch Anwendung der bedingten Markoffschen Zufallsfelder sinkt die Klassifikationsgenauigkeit der Klasse Gesund dagegen auf komplett gesunden leicht und auf mit *Cercospora beticola* infizierten Blättern etwas stärker ab. Auf dem völlig gesunden Blatt in Tabelle 4.11 findet eine Reduktion der Klassifikationsgenauigkeit von 98.62% auf 98.55% statt und auf dem mit *Cercospora beticola* infizierten Blatt von 86.24% auf 84.62% statt. Auf dem mit *Uromyces betae* infizierten Blatt steigt sie dagegen von 98.73% auf 98.84%. Im Median bedeutet dies eine Abnahme von 96.65% auf 96.40% (siehe Tabellen 4.6 und 4.8). Der stärkste Rückgang besteht damit auf dem mit *Cercospora beticola* befallenen Blatt. Dies ist darauf zurückzuführen, dass die Anwendung der Markoffschen Zufallsfelder zu einer Vergrößerung der detektierten *Cercospora*-Blattflecken führt und dadurch auch von den Ground Truth Daten her gesunde Pixel als krank deklariert werden. Als *Cercospora beticola* fehlklassifizierte Areale in der Bayesklassifikation werden ebenfalls durch Anwendung der bedingten Markoffschen Zufallsfelder vergrößert, so dass dieser Effekt auch auf *Uromyces*-Bildern und Bildern von gesunden Blättern beobachtet werden kann.

Der große Vorteil der bedingten Markoffschen Zufallsfelder liegt aber in der effektiven Eliminierung der fehlklassifizierten Braunrostpixel am Rand der *Cercospora*-Flecken. Dieses Phänomen ist sowohl in den Ausschnitten in Abbildung 4.8 gut zu sehen als auch an der Fehlklassifikationsrate in Tabelle 4.9 erkennbar, welche von 8.21% auf 0.29% absinkt. Desweiteren kann durch die Glättung die Fehlklassifikation der Blattadern als *Uromyces betae* vor allen Dingen bei Anwendung der bedingten Markoffschen Zufalls-

felder deutlich reduziert werden, wie in Abbildung 4.10 dargestellt ist. Gleichzeitig ist dabei die Detektion korrekter Braunrost-Pusteln gut, wie in Abbildung 4.9 visualisiert ist.

4.4 Regionenbasierte Klassifikation

Nach Glättung des pixelweisen Ergebnisses der adaptiven Bayesklassifikation existieren immer noch eine relativ große Anzahl an falschen Blattflecken, wie z.B. als krank klassifizierte Blattadern, oder Blattflecken, welche teilweise der jeweils anderen Blattkrankheit zugeordnet wurden (siehe Abbildungen 4.8, 4.9 und 4.10). Diese Fehler sollen nun durch eine regionenbasierte Klassifikation eliminiert werden.

Die für eine regionenbasierte Klassifikation erforderliche Segmentierung des Bildes ist in diesem Falle durch die zuvor durchgeführte pixelweise Bayesklassifikation und die anschließende Glättung erfolgt. Prinzipiell besteht aber auch die Möglichkeit, das Bild z.B. mit dem Wasserscheidenalgorithmus zu segmentieren. Diese Art der Segmentierung zeigte aber bei den hier durchgeführten Untersuchungen keinen Erfolg, wie dem folgenden, ersten Unterkapitel entnommen werden kann. Das zweite Unterkapitel beschreibt dann die Durchführung der regionenbasierten Klassifikation anhand der geglätteten Labelbilder und das letzte Unterkapitel präsentiert schließlich die Ergebnisse der regionenbasierten Klassifikation.

4.4.1 Einsatz von Wasserscheiden zur Regionenbildung

Dieses Kapitel untersucht inwieweit der Wasserscheidenalgorithmus zur Regionenbildung bei Blattkrankheiten geeignet ist. Denn statt einer pixelweisen Klassifikation mit anschließender Glättung kann ein Bild auch direkt mit Hilfe von Segmentierungsverfahren, wie dem Wasserscheidenalgorithmus, in Regionen aufgeteilt werden. Der im Folgenden angewandte Wasserscheidenalgorithmus stammt von Vincent und Soille (1991).

Zur Segmentierung der Bilder mit dem Wasserscheidenalgorithmus wurden zwei Experimente durchgeführt. Im ersten Experiment erfolgte die Anwendung des Wasserscheidenalgorithmus' auf RGB-Bildern und im zweiten Experiment auf dem Infrarotkanal von MS-Bildern.

Wie in den Abbildungen 4.11 bis 4.15 zu sehen ist, ist es eher schwierig, eine gute nicht-semantische Vorsegmentierung mit Hilfe des Wasserscheidenalgorithmus' zu erhalten. Die erste Abbildung 4.11 zeigt das Ergebnis des Wasserscheidenalgorithmus' auf dem RGB-Bild eines mit *Cercospora beticola* infizierten Zuckerrübenblattes. Wie dort dargestellt ist, hat der Wasserscheidenalgorithmus sehr häufig die Grenze zwischen dem braunen Rand des *Cercospora*-Blattflecks und der gesunden grünen Blattfläche nicht erkannt. Der Kontrast zwischen diesen Bereichen ist offensichtlich nicht hoch genug, da der Übergang zwischen gesunder und kranker Blattfläche oft fließend ist. Zur Ver-



Abbildung 4.11: Wasserschneiden auf dem RGB-Bild eines mit *Cercospora beticola* infizierten Zuckerrübenblattes. Die Grenze zwischen gesunden und kranken Blattbereichen wurde fast nie gefunden, da der Kontrast zwischen dem dunkelbraunen Rand eines Blattfleckes und der gesunden, grünen Blattfläche zu gering ist. Beispiele für nicht korrekt segmentierte Blattfleckenränder sind gelb markiert.

deutlichung ist in Abbildung 4.12 die Vergrößerung eines Blattflecks zu sehen, wo die korrekte Grenze mit weißer Farbe eingezeichnet ist.

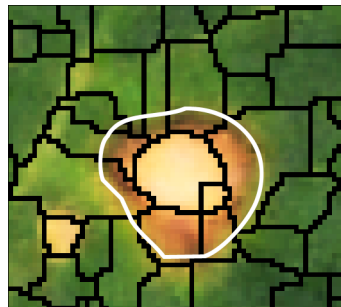


Abbildung 4.12: Diese Abbildung zeigt die Vergrößerung eines einzelnen Blattflecks. Der Wasserscheidenalgorithmus hat die Grenze um den *Cercospora*-Fleck nicht gefunden. Der korrekte Übergang von kranker, brauner zu gesunder, grüner Blattfläche ist durch den weißen Ring angedeutet.

Die Anwendung des Wasserscheidenalgorithmus' auf einem mit *Uromyces betae* infizierten Blattes blieb in diesem Experiment völlig erfolglos, wie in Abbildung 4.13 zu sehen ist. Trotz guten Kontrastes zwischen gesunder Blattfläche und den Rostpusteln konnten die Pusteln aufgrund ihrer sehr geringen Größe nicht segmentiert werden. Evt. könnte die Segmentierung der Pusteln durch eine geringere Rauschunterdrückung verbessert werden, damit würde aber gleichzeitig auch die Anzahl der Wasserschneiden, welche nur auf Rauschen basieren, deutlich ansteigen. Eine andere Möglichkeit wäre

es, die Bilder in einer deutlich höheren Auflösung aufzunehmen. Dies dürfte aber im Falle eines Einsatzes von dem Verfahren im Feld eher schwierig werden.



Abbildung 4.13: Die kleinen Uromyces Pusteln wurden wegen der geringen Größe vom Wasserscheidenalgorithmus überhaupt nicht gefunden, obwohl an sich der Kontrast zwischen Blatt und Pustel sehr gut ist.

Die Idee im zweiten Experiment war es, aufgrund des besseren Kontrastes zwischen kranker und gesunder Blattfläche im Infrarotbereich den Wasserscheidenalgorithmus auf dem Infrarotkanal des Multispektralbildes anzuwenden. Wie in Abbildung 4.14 zu sehen ist, funktioniert dieses Vorgehen auf einem mit *Cercospora beticola* infizierten Blatt, dessen gesunde Blattfläche noch viel Chlorophyll enthält, hervorragend. Die Segmentierung erfolgte anhand des mittleren Blattes in dieser Abbildung. Zur Anschauung sind die Wasserscheiden zusätzlich links im entsprechenden Ausschnitt des RGB-Bildes und rechts des jeweiligen MS-Bildes eingezeichnet.

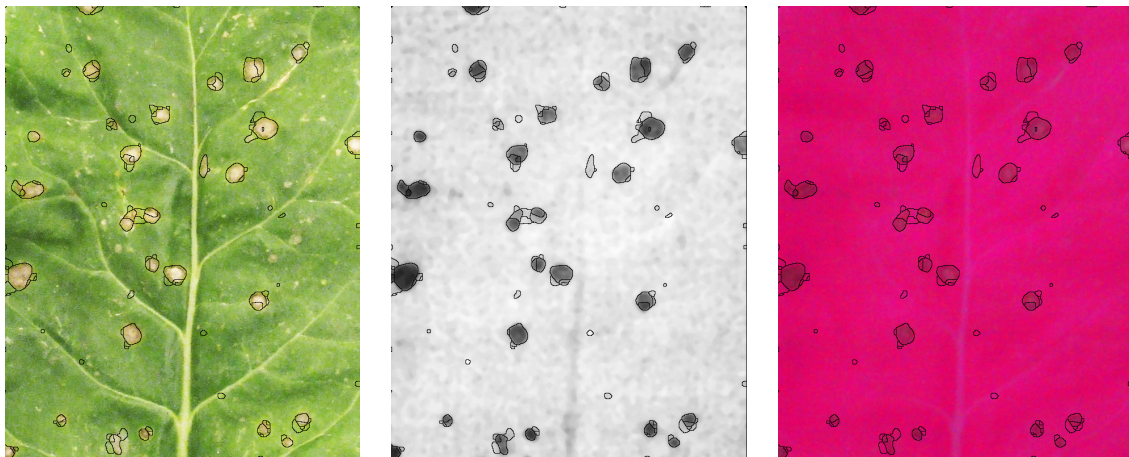


Abbildung 4.14: Hervorragendes Ergebnis des Wasserscheidenalgorithmus'. Alle Blattflecken wurden richtig segmentiert. Links sind die Wasserscheiden im RGB-Bild eingezeichnet. In der Mitte ist das Grauwertbild abgebildet, anhand dessen die Wasserscheiden berechnet wurden und rechts ist das MS-Bild mit eingezeichneten Wasserscheiden zu sehen.

4 Experimente

Sobald allerdings der Befall des Blattes so stark wird, dass eine Schädigung des kompletten Blattes eintritt, treten auch bei diesem Vorgehen Probleme auf, wie in Abbildung 4.15 zu sehen ist. Durch den Abbau des Chlorophylls in der ehemals gesunden Blattfläche besteht kein Kontrast mehr zwischen gesunder und kranker Blattfläche im Infrarotkanal. Daher konnten kaum Regionen segmentiert werden und eine Erkennung der Grenzen der Blattflecken erfolgte daher nicht.

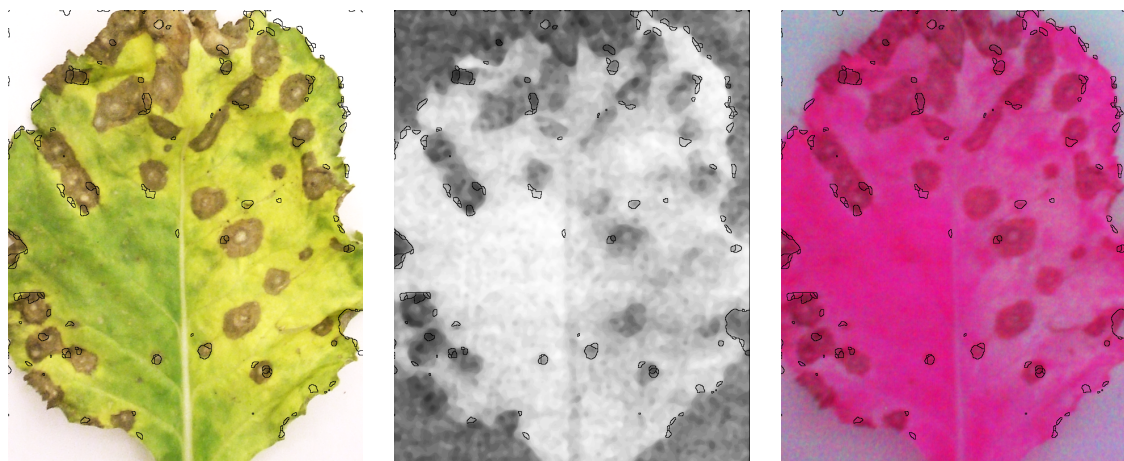


Abbildung 4.15: Der Befall des Blattes mit *Cercospora beticola* ist schon sehr weit fortgeschritten. Dadurch ist das Blatt insgesamt sehr angegriffen und damit auch die Photosyntheseleistung der noch nicht mit Blattflecken bedeckten Blattfläche herabgesetzt. Dies führt dazu, dass im Infrarotkanal kaum mehr ein Unterschied zwischen Blattflecken und Blattfläche ohne Flecken besteht. Genauso wie in der vorherigen Abbildung sind links die Wasserscheiden im RGB-Bild eingezeichnet. In der Mitte ist das Grauwertbild abgebildet anhand dessen die Wasserscheiden berechnet wurden. Und rechts ist das MS-Bild mit eingezeichneten Wasserscheiden dargestellt.

Die Segmentierung der Blattkrankheiten mit dem Wasserscheidenalgorithmus führte damit in diesen Experimenten zu eher unbefriedigenden Ergebnissen und wurde daher nicht weiter eingesetzt.

4.4.2 Maximum-Likelihood Klassifikation der geglätteten Regionen

Die im Folgenden beschriebene regionenbasierte Klassifikation basiert auf den geglätteten Labelbildern der pixelweisen, adaptiven Bayesklassifikation. Die Durchführung der ML-Klassifikation der Regionen ist in Kapitel 3.3.6 erläutert. Es wurden sowohl die mittels des Majoritätsfilters als auch die anhand der bedingten Markoffschen Zufallsfelder geglätteten Bilder klassifiziert. Wie in Kapitel 4.3.1 dargelegt, wurden 180 Bilder mit

dem Majoritätsfilter geglättet und aufgrund der sehr langen Berechnungszeiten nur 12 der 180 Bilder mit den bedingten Markoffschen Zufallsfeldern.

Zum Training des ML-Klassifikators wurden aus dem Datensatz von 180 Bildern die 12 mit den bedingten Markoffschen Zufallsfeldern geglätteten Bilder entfernt. Aus dem resultierenden Datensatz von 168 Bildern wurden dann 45 Bilder zufällig ausgewählt. Die restlichen 123 gefilterten Bilder sowie alle mittels des CRFs geglätteten Bilder wurden anschließend klassifiziert.

Die Klassifikation erfolgte in zwei Schritten. Im ersten Schritt wurden die Regionen anhand ihrer mittleren Rot-, Grün-, Blau- und Infrarot-Farbwerte in die Klassen Gesund und Krank aufgeteilt. Der zweite Schritt diente zur Zuordnung der im vorigen Schritt als krank deklarierten Regionen zu den beiden Blattkrankheiten *Cercospora beticola* und *Uromyces betae*. Diese Zuordnung basierte auf der Fläche, des Umfangs und der Rundheit der jeweiligen Region.

4.4.3 Ergebnis der regionenbasierten ML-Klassifikation

Die Tabellen 4.12 bis 4.16 enthalten die Ergebnisse der regionenbasierten Klassifikation. Die beiden ersten Tabellen 4.12 und 4.13 zeigen dabei die minimal, im Median und maximal erzielten Klassifikationsgenauigkeiten für die im vorigen Kapitel 4.3.2 ausgewählten zwölf Blätter. Die folgenden drei Tabellen 4.14, 4.15 und 4.16 zeigen die Ergebnisse der regionenbasierten Klassifikation auf den ebenfalls im vorigen Kapitel ausgewählten einzelnen Blättern. In der Tabelle 4.14 ist somit das Ergebnis für das Blatt mit der Blattfleckenkrankheit *Cercospora beticola* dargestellt, in Tabelle 4.15 für das mit Braunrost befallene Blatt und in Tabelle 4.16 schließlich für das gesunde Blatt. In diesen drei Tabellen steht in der ersten und dritten Zeile jeweils zum Vergleich das Ergebnis der Majoritätsfilterung bzw. der Glättung mit den bedingten Markoffschen Zufallsfeldern (CRF). Der zweiten Zeile kann das Ergebnis der Regionen-Klassifikation basierend auf dem Resultat der Majoritätsfilterung und der letzten Zeile die Regionen-Klassifikation basierend auf dem CRF-Ergebnis entnommen werden. Alle Ergebnisse der regionenbasierten Klassifikation zeigen das endgültige Resultat nach der Farb- und Form-Klassifikation. Zum besseren Vergleich mit den vorigen Ergebnissen wurden bei der regionenbasierten Klassifikation zwar Regionen klassifiziert, das Ergebnis der Klassifikation aber ebenfalls auf Pixelebene dargestellt.

Wie den Tabellen entnommen werden kann, erhöht sich in allen Fällen die Klassifikationsgenauigkeit der Klasse Gesund. So steigt sie bei der Regionenklassifikation basierend auf dem Filterungsergebnis im Median von 97% auf fast 100% (siehe Tabellen 4.7 und 4.12).

Ebenfalls gesteigert wird fast immer die Klassifikationsgenauigkeit der Blattfleckenkrankheit *Cercospora beticola*. Bei der Regionenklassifikation auf dem Filterungsergebnis nimmt hier die Klassifikationsgenauigkeit im Median von 85% auf 88% zu (siehe Tabellen 4.7 und 4.12). Die höchste Klassifikationsgenauigkeit für diese Klasse kann mit

Regionenklassifikation I			
%	$\hat{h}$	$\hat{i}_{\text{cerc}}$	$\hat{i}_{\text{urom}}$
$\tilde{h}$	90.76	0.00	0.00
	99.75	0.25	0.05
	100.00	9.08	0.16
$\tilde{i}_{\text{cerc}}$	11.43	78.86	0.00
	11.84	88.16	0.00
	12.23	88.57	8.91
$\tilde{i}_{\text{urom}}$	25.00	0.00	20.29
	46.38	16.10	26.30
	57.60	33.33	75.00

Tabelle 4.12: Konfusionsmatrix für die Regionenklassifikation basierend auf dem Filterungsergebnis aus Tabelle 4.7.

Regionenklassifikation II			
%	$\hat{h}$	$\hat{i}_{\text{cerc}}$	$\hat{i}_{\text{urom}}$
$\tilde{h}$	86.17	0.00	0.00
	99.61	0.36	0.09
	100.00	13.52	0.32
$\tilde{i}_{\text{cerc}}$	5.73	83.85	0.00
	7.01	90.99	0.42
	8.59	94.27	9.14
$\tilde{i}_{\text{urom}}$	17.86	0.00	38.78
	37.68	0.00	62.32
	48.41	12.81	82.14

Tabelle 4.13: Konfusionsmatrix für die Regionenklassifikation basierend auf den bedingten Markoffschen Zufallsfeldern aus Tabelle 4.6.

Regionenbasierte Klassifikation				
Blattfleckenkrankheit				
%		$\hat{h}$	$\hat{i}_{\text{cerc}}$	$\hat{i}_{\text{urom}}$
$\tilde{h}$	Filterung	87.44	9.46	3.10
	Filterung Regionen	90.76	9.08	0.16
	CRF	84.62	14.26	1.13
	CRF Regionen	86.17	13.52	0.32
$\tilde{i}_{\text{cerc}}$	Filterung	8.75	85.21	6.04
	Filterung Regionen	11.84	88.16	0
	CRF	5.73	93.98	0.29
	CRF Regionen	5.73	94.27	0

Tabelle 4.14: Vergleich der Klassifikationsgenauigkeiten auf einem mit der Blattfleckenkrankheit *Cercospora beticola* befallenen Blatt. In der ersten Zeile steht das Ergebnis der Glättung mit dem Majoritätsfilter und in der folgenden Zeile das Ergebnis für die Regionenklassifikation des gefilterten Bildes. In der dritten Zeile kommt das Ergebnis für die Glättung mit den bedingten Markoffschen Zufallsfeldern und in der letzten Zeile steht schließlich das Resultat der Regionenklassifikation basierend auf dem CRF-Ergebnis. Alle Genauigkeiten wurden auf Pixelebene berechnet.

Regionenbasierte Klassifikation				
Braunrost				
%		$\hat{h}$	$\hat{i}_{\text{cerc}}$	$\hat{i}_{\text{urom}}$
$\tilde{h}$	Filterung	98.98	0.72	0.30
	Filterung Regionen	99.87	0.06	0.07
	CRF	98.84	0.96	0.21
	CRF Regionen	99.79	0.08	0.13
$\tilde{i}_{\text{urom}}$	Filterung	7.82	3.85	88.32
	Filterung Regionen	57.60	16.10	26.30
	CRF	7.71	9.41	82.88
	CRF Regionen	48.41	12.81	38.78

Tabelle 4.15: Der Aufbau der Tabelle ist genauso wie in der vorigen. Die Ergebnisse in dieser Tabelle beziehen sich in diesem Fall aber auf ein mit dem Braunrost *Uromyces betae* befallenes Blatt.

Regionenbasierte Klassifikation				
Gesund				
%		$\hat{h}$	$\hat{i}_{\text{cerc}}$	$\hat{i}_{\text{urom}}$
$\tilde{h}$	Filterung	98.75	1.20	0.05
	Filterung Regionen	99.78	0.21	0.01
	CRF	98.55	1.44	0.01
	CRF Regionen	99.73	0.26	0.01

Tabelle 4.16: Die Ergebnisse in dieser Tabelle beziehen sich auf ein gesundes Blatt.

94% für ein stark mit dieser Krankheit infiziertes Blatt bei einer Regionenklassifikation basierend auf dem CRF-Ergebnis erzielt werden (siehe Tabelle 4.14). Im Median sinkt allerdings bei der Regionenklassifikation basierend auf dem CRF-Ergebnis die Genauigkeit für diese Klasse von 93% auf 91% (siehe Tabellen 4.6 und 4.13).

Die Klassifikationsgenauigkeit des Braunrostes *Uromyces betae* sinkt im Median ganz erheblich auf 26% bzw. 62% ab (siehe Tabellen 4.12 und 4.13). Bei dem stark mit Braunrost befallenen Blatt sinkt die Klassifikationsgenauigkeit bei der Regionenklassifikation basierend auf dem CRF-Ergebnis von 83% auf 39% ab (siehe Tabelle 4.15). Wie schon in Kapitel 3.3.3 angedeutet, ist dies voraussichtlich darauf zurückzuführen, dass die vorab als Braunrost segmentierten Regionen größer sind als die zugehörigen Ground Truth Regionen. Dadurch beinhaltet die segmentierte Braunrostregion einen hohen Anteil gesunder Blattfläche und wird dementsprechend fehlklassifiziert.

4.5 Untersuchung der Übertragbarkeit des Verfahrens

Alle bisher durchgeführten Experimente basierten auf den in Kapitel 3.3.2 vorgestellten Bildern von Zuckerrübenblättern mit zwei verschiedenen Blattkrankheiten. Es existieren aber eine sehr große Anzahl weiterer Nutzpflanzen und unterschiedlichster Faktoren, welche zu einem Absterben der Blätter bzw. der Pflanzen führen können. Es wäre daher von großem Interesse, dass das hier entwickelte Verfahren auch in diesen Fällen zur Klassifikation von Blattläsionen eingesetzt werden könnte.

Das folgende Experiment basiert auf Reisblättern. Im Gegensatz zu den bisher durchgeführten Experimenten mit biotischen Stressfaktoren wurden die Reispflanzen einer hohen Konzentration an Eisen ausgesetzt, d.h. also einem abiotischen Stressfaktor. Die daraus resultierende Eisentoxizität ist bei den Reispflanzen eine der wichtigsten Stressreaktionen im Nassreisanbau (Asch, 2007). Das Ziel dieses Experimentes ist die Klassifikation der Symptome der Eisentoxizität auf den Blättern der Reispflanzen. Als Symptom der Eisentoxizität treten braun-gelbe Verfärbungen der Blätter auf, welche sich entlang der Blattfasern ausbreiten (siehe Abbildungen 4.18 und 4.19).

4.5.1 Datenmaterial

Zur Untersuchung der automatischen Erkennbarkeit der Eisentoxizität auf Reisblättern wurden Aufnahmen von 72 Blättern von insgesamt sechs verschiedenen Reissorten gemacht. Drei der Reissorten reagieren eher sensitiv auf die Zuführung hoher Eisengaben und die anderen drei Reissorten gelten als eher resistent. Mit der automatischen Erkennung der befallenen Blattfläche könnte hier also eine objektive Aussage über die Sensitivität verschiedener Sorten getroffen werden.

Für jede Reissorten standen 12 Pflanzen zur Verfügung. Jeweils der Hälfte der Pflanzen wurde Eisen zugeführt, um eine Stressreaktion zu induzieren. Die andere Hälfte diente zur Kontrolle und wurde demnach möglichst optimal gehalten. Von jeder Reis-

pflanze wurde das jüngste Fahnenblatt fotografiert. Denn sobald selbst das jüngste Fahnenblatt Symptome zeigt, stirbt die Pflanze mit hoher Wahrscheinlichkeit komplett ab.

Die Aufnahme der Reisblätter erfolgte genauso wie die der Zuckerrübenblätter im Labor unter kontrollierter Beleuchtung. Die Bilder wurden mit der handelsüblichen RGB-Kamera Nikon D200 mit einer Auflösung von 3872×2592 Pixeln aufgenommen. Bei der vorliegenden Aufnahmehöhe von ungefähr 57 cm entspricht 1 Pixel damit ca. 0.066 mm im Quadrat. Auf einem Bild sind dabei jeweils alle sechs Blätter einer Varietät abgebildet (siehe z.B. Abbildung 4.16).

4.5.2 Durchführung der Klassifikation

Die Klassifikation der Reisblätter erfolgte sowohl manuell als auch automatisch. Die manuelle Klassifikation wurde dabei durch einen Experten direkt am Aufnahmetag anhand der realen Blätter durch subjektive Zuweisung eines Wertes zwischen null und zehn durchgeführt. Wie in Tabelle 4.17 verdeutlicht ist, kann vom dem zugewiesenen Wert direkt auf die subjektiv geschätzte, befallene Blattfläche geschlossen werden. Ein Wert von 4 bedeutet z.B., dass 40% der Blattfläche bereits zerstört ist. Die Zuweisung des Wertes erfolgte dabei in 5%-Schritten.

Das Experiment zur automatischen Klassifikation der Reisblätter diente dazu einen ersten Anhaltspunkt zur Erkennbarkeit der Eisentoxizität auf Reisblättern zu erhalten. Zu diesem Zwecke wurde eine pixelweise ML-Klassifikation eingesetzt. Denn wie in Kapitel 4.2.5 herausgefunden wurde, können mit einer ML-Klassifikation bessere Klassifikationsgenauigkeiten als mit einer MAP-Klassifikation erzielt werden. Die Genauigkeit der ML-Klassifikation könnte dann zwar noch durch eine optimierte Gewichtungsfunktion bei Einsatz einer Bayesklassifikation verbessert werden, für einen ersten Anhaltspunkt genügt aber auch die ML-Klassifikation. Entsprechend wurde in diesem Experiment für die beiden Klassen, gesunde und zerstörte Blattfläche, jeweils eine Gaußsche Mischverteilung aus zwei Verteilungen angenommen. Die Klassifikation erfolgte anhand der Farbinformationen der Kanäle Rot, Grün und Blau.

Zum Training des Klassifikators wurden in einem Bild, auf den dort abgebildeten sechs Blättern, einige Eisentoxizitätsstellen manuell gelabelt und anhand dessen die Parameter der Mischverteilung berechnet. Die gesunde Blattfläche wurde für jede Sorte anhand des ersten Kontrollbildes trainiert, da sich in Voruntersuchungen gezeigt hat, dass die Grünfärbung von Sorte zu Sorte variiert.

4.5.3 Ergebnis der Klassifikation

In den Abbildungen 4.16, 4.18 und 4.19 sind beispielhaft die Klassifikationsergebnisse für drei Bilder dargestellt. Alle drei Abbildungen zeigen auf der linken Seite das jeweilige RGB-Bild, auf welchem die Klassifikation basiert, und auf der rechten Seite das zugehörige Labelbild der ML-Klassifikation. In den Labelbildern repräsentiert weiß

Symptomstärke	
Befallene Blattfläche	Zugewiesener Wert
0%	0
10%	1
20%	2
30%	3
40%	4
50%	5
60%	6
70%	7
80%	8
90%	9
100%	10

Tabelle 4.17: Die Blätter werden nach dem Grad ihres Befalls mit Eisentoxizität manuell einem Wert zwischen 0 und 10 zugeordnet.

4 Experimente

den Hintergrund, grau die gesunde Blattfläche und schwarz die durch die Eisentoxizität zerstörte Blattfläche.

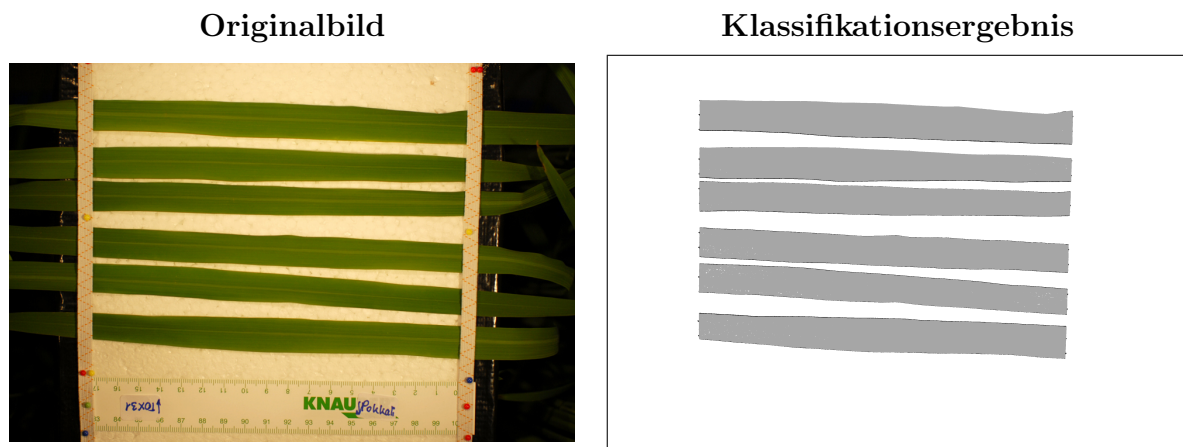


Abbildung 4.16: Automatisches Klassifikationsergebnis der Blätter der Sorte *aTOX3107* am ersten Tag. Die Blätter weisen noch keinerlei Symptome auf. Der manuell zugewiesene Wert beträgt dementsprechend 0. Die automatisch als befallene Blattfläche berechnete Fläche beträgt 1.58 % durch die fehlklassifizierten Blattränder.

Die erste Abbildung 4.16 zeigt das Klassifikationsergebnis für ein gesundes Blatt. Entsprechend wurde dem Bild manuell ein Wert von 0 zugewiesen. Bei der automatischen Klassifikation wurden 1.58% der gesamten Blattfläche als durch Eisentoxizität zerstört deklariert. Dies ist auf die Fehlklassifikation der Blattränder zurückzuführen, wie in Abbildung 4.17 in einer Vergrößerung dargestellt ist. Die Fehlklassifikation basiert auf dem im RGB-Bild braun erscheinenden Blattrand. Der braune Rand ist in der Realität nicht vorhanden und ist höchstwahrscheinlich auf Schattenwurf zurückzuführen. Durch Wahl einer anderen Aufnahmegeometrie bzw. anderer Hintergründe dürfte dieses Problem aber hinfällig werden.

Die Abbildungen 4.18 und 4.19 zeigen die Stressreaktionen am sechsten Tag nach Beginn des Versuchs bei der Sorte *aTOX3107* und der Sorte *Nipponbare*. In beiden Fällen weicht die manuelle Klassifikation, welche auf dem subjektiven Eindruck eines Experten basiert, und der automatisch detektierte Wert nur geringfügig voneinander ab. Es ist somit kein signifikanter Unterschied nachweisbar.

Die automatische Klassifikation der Eisentoxizität auf Reisblättern bietet damit z.B. die Möglichkeit objektiv die Sensitivität verschiedener Sorten im Hinblick auf die Eisentoxizität zu untersuchen. Desweiteren zeigt die erfolgreiche Anwendung des Verfahrens zur Klassifikation eines abiotischen Stressfaktors auf Reisblättern, dass das Verfahren potentiell auch zur Detektion von Blattläsionen auf anderen Blättern eingesetzt werden kann.

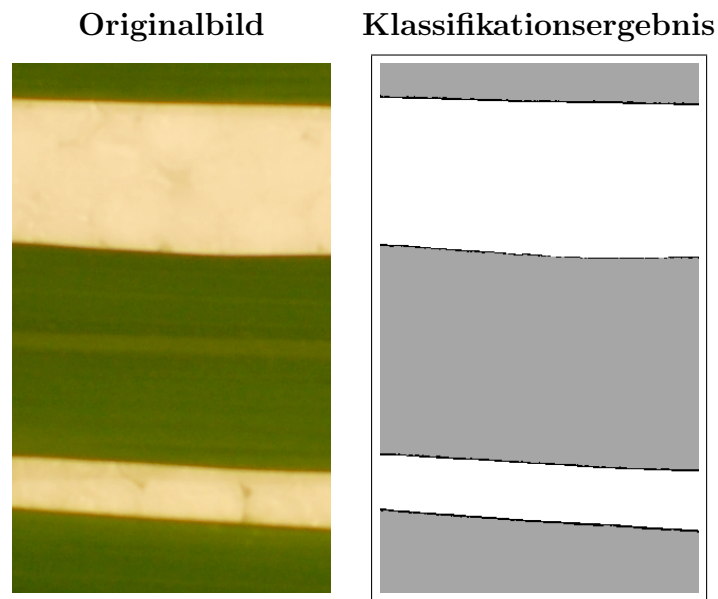


Abbildung 4.17: Blattrandproblematik: Der Blattrand erscheint im Bild braun, so dass er als befallene Fläche klassifiziert wird.

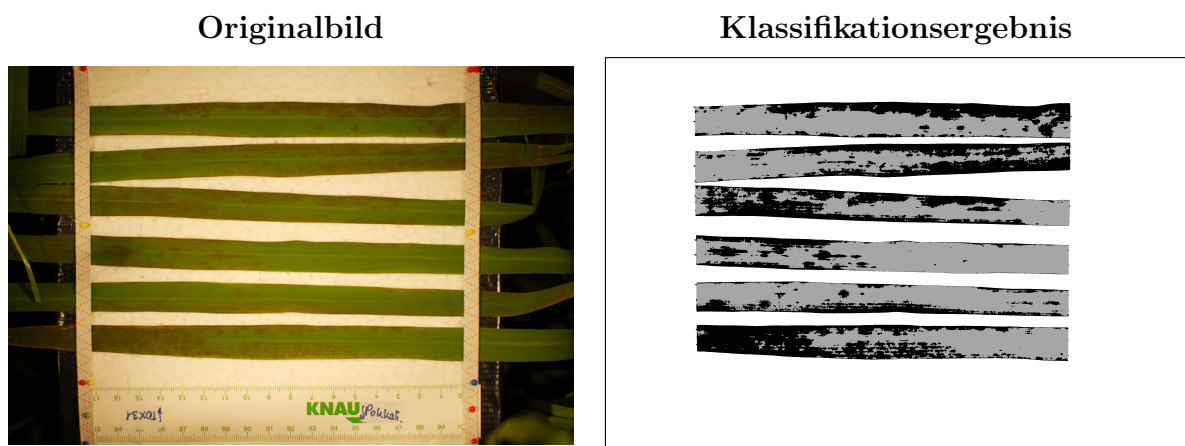


Abbildung 4.18: Automatisches Klassifikationsergebnis der Blätter der Sorte *aTOX3107* am sechsten Tag. Die Blätter weisen starke Symptome auf. Der manuell zugewiesene Symptomwert beträgt 3.5, d.h. es wurde geschätzt, dass 35% der Blattfläche zerstört sind. Die automatisch als befallene Blattfläche berechnete Fläche beträgt 38.88 %. Der manuell zugewiesene und der automatisch geschätzte Wert unterscheiden sich damit nicht signifikant.

Originalbild



Klassifikationsergebnis

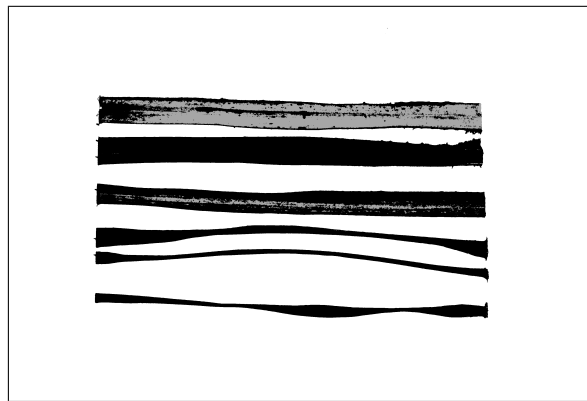


Abbildung 4.19: Automatisches Klassifikationsergebnis der Blätter der Sorte *Nipponbare* am sechsten Tag. Die Blätter weisen starke Symptome auf bzw. sind bereits abgestorben. Der manuell zugewiesene Symptomwert beträgt hier 7.75, d.h. es wurde geschätzt, dass 77.5% der Blattfläche zerstört sind. Die automatisch als befallene Blattfläche berechnet Fläche beträgt 80.41 %. Auch in diesem Fall ist kein signifikanter Unterschied zwischen den beiden Werten feststellbar.

5 Schlussfolgerung

Bei der Entwicklung des Verfahrens zur automatischen Detektion von Krankheiten auf Blättern von Nutzpflanzen erfolgte die Verifikation von vier globalen Arbeitshypothesen. Laut der ersten Arbeitshypothese erhöhen sich die Klassifikationsgenauigkeiten aller Klassen signifikant bei einer Berücksichtigung der 4-er Nachbarschaft in den Merkmalen. Wie die Ergebnisse dieser Arbeit zeigen, ist dies auch tatsächlich der Fall. So konnte die Klassifikationsgenauigkeit der gesunden Blattfläche um 7% von 84% auf 91% gesteigert werden, die Klassifikationsgenauigkeit der Blattfleckenkrankheit ebenfalls um 7% von 74% auf 81% und die Klassifikationsgenauigkeit des Braunrostes sogar um 17% von 59% auf 76% (siehe Kapitel 4.2.2).

Die nächste Arbeitshypothese besagte, dass die Klassifikationsgenauigkeit des seltenen Braunrostes durch eine Optimierung des Bayesklassifikators weiter erhöht werden kann. Dies traf auch zu, denn die Klassifikationsgenauigkeit des Braunrostes schwankte z.B. allein durch die gewählte Anzahl der Gaußschen Verteilungen in der Likelihood-Funktion zwischen 63% und 81% (siehe Kapitel 4.2.4). Durch die Wahl der prior-Kostenmatrix (siehe Kapitel 4.2.6) in der Bayesklassifikation konnten sogar 95% des Braunrostes korrekt klassifiziert werden. Dafür betrug hier aber die Klassifikationsgenauigkeit der gesunden Blattfläche nur 89% und die Klassifikationsgenauigkeit der Blattfleckenkrankheit lediglich 86%. Die hohe Klassifikationsgenauigkeit des Braunrostes wurde also durch niedrigere Klassifikationsgenauigkeiten bei den anderen beiden Klassen erkauft. Um die Fehlklassifikationsraten dieser beiden zahlenmäßig deutlich häufiger auftretenden Klassen zu minimieren, wurde vorgeschlagen, die posterior-Kostenmatrix zu verwenden. Denn dort beträgt die Klassifikationsgenauigkeit des Braunrostes immerhin noch 86% und die Klassifikationsgenauigkeit der gesunden Blattfläche dafür 94% und die der Blattfleckenkrankheit 91%.

Mit der darauffolgenden Arbeitshypothese sollte die Behauptung überprüft werden, ob die Glättung des pixelweisen Klassifikationsergebnisses mit bedingten Markoffschen Zufallsfeldern einer Glättung durch einen Majoritätsfilter überlegen ist. In Bezug auf die Klassifikationsgenauigkeit der Blattfleckenkrankheit konnte dies auch bestätigt werden, da sie bei einer Glättung durch die Zufallsfelder im Median bei 93% liegt und damit um 8% höher als bei einer Majoritätsfilterung. Sowohl die Klassifikationsgenauigkeit der gesunden Blattfläche als auch die des Braunrostes sind allerdings bei der Majoritätsfilterung höher. So kann z.B. für den Braunrost durch die Filterung im Median eine Genauigkeit von 88% erzielt werden und mit den Zufallsfeldern nur eine Genauigkeit von 83% (siehe Kapitel 4.3.2). Unabhängig von den erzielten Klassi-

fikationsgenauigkeiten liegt der große Vorteil der bedingten Markoffschen Zufallsfelder aber in der effektiven Eliminierung der fehlklassifizierten Braunrostpixel am Rand der *Cercospora*-Flecken, wie in Abbildung 4.8 auf Seite 76 zu sehen ist. Ebenfalls reduziert die Anwendung der bedingten Markoffschen Zufallsfelder deutlich stärker als die Majoritätsfilterung die Fehlklassifikation von Blattadern als *Uromyces betae*, wie in Abbildung 4.10 auf Seite 78 visualisiert ist. Die bedingten Markoffschen Zufallsfelder zeigen damit großes Potential in der Klassifikation von Blattkrankheiten.

In der letzten Arbeitshypothese wurde die Annahme aufgestellt, dass eine Regionenklassifikation basierend auf den zuvor geglätteten Bildern die Klassifikationsgenauigkeiten aller Klassen erhöht. Auf die Klassifikationsgenauigkeit der gesunden Blattfläche trifft dies auch zu. So konnte für die gesunde Blattfläche bei beiden Regionenklassifikationen Genauigkeiten von über 99% erzielt werden, während die Genauigkeit bei der Filterung bei 97% und bei den Markoffschen Zufallsfeldern bei 96% lag. Die Klassifikationsgenauigkeit der Blattfleckenkrankheit nahm bei der Regionenklassifikation basierend auf dem Filterungsergebnis von 85% auf 88% zu, bei der Regionenklassifikation basierend auf den Markoffschen Zufallsfeldern dagegen von 93% auf 91% ab. Ganz drastisch sank die Klassifikationsgenauigkeit des Braunrostes bei den beiden Regionenklassifikationen. Sie beträgt dort basierend auf den Markoffschen Zufallsfeldern 62% und basierend auf dem Filterungsergebnis sogar nur 26% (siehe Kapitel 4.4.3). Zur Verminderung von Fehlklassifikationen gesunder Blattareale ist die hier durchgeführte Regionenklassifikation somit gut geeignet, zur besseren Klassifikation der beiden Blattkrankheiten, insbesondere des Braunrostes, müssen aber noch weitergehende Forschungsarbeiten durchgeführt werden. Evt. kann die Regionenklassifikation bei einer weiteren Optimierung der Markoffschen Zufallsfelder auch ganz entfallen.

Zusammenfassend lässt sich sagen, dass das hier vorgestellte Verfahren zur Detektion von Blattkrankheiten großes Potential bietet. Die pixelweise, adaptive Bayesklassifikation dürfte sich bei Vorhandensein von Trainingsdaten leicht auf die Erkennung von Krankheiten auf anderen Blättern anpassen lassen. Denn die meisten Parameter werden im Trainingsprozess automatisch gelernt, lediglich die aus der Verteilung der Daten resultierende Anzahl an Clustern in der Mischverteilung wurde für jede Klasse in dieser Arbeit experimentell ermittelt. Die darauffolgende Majoritätsfilterung benötigt ebenfalls keinerlei Anpassung an andere Blätter oder Krankheiten. Die alternativ durchgeführte Glättung mit den bedingten Markoffschen Zufallsfeldern setzt momentan allerdings noch eine experimentelle Ermittlung der Matrix T voraus. Die abschließend durchgeführte Regionenklassifikation funktioniert wieder vollautomatisch.

Die in dieser Arbeit durchgeführte Untersuchung zur Übertragbarkeit des Verfahrens zeigte, dass das Verfahren durchaus auch zur Klassifikation von abiotischen Stressfaktoren auf den Blättern anderer Nutzpflanzen, wie in diesem Fall der Eisentoxizität auf Reisblättern, eingesetzt werden kann. Die visuelle Überprüfung der Klassifikationsergebnisse auf den Reisblättern weist sogar auf höhere Klassifikationsgenauigkeiten als bei den Zuckerrübenblättern hin (siehe Kapitel 4.5). Dies dürfte auf die einheitli-

chere Struktur der Reisblätter zurückzuführen sein, welche keine Blattadern wie die Zuckerrübenblätter aufweisen.

Mit dem in dieser Arbeit vorgestellten Klassifikationsverfahren zur automatischen Detektion von Krankheiten auf Blättern von Nutzpflanzen besteht somit die Möglichkeit in kurzer Zeit große Datensätze auf das Vorhandensein von Blattkrankheiten sowie auf die Verteilung der Blattkrankheiten auf den Blättern zu überprüfen. Das Klassifikationsergebnis ist dabei objektiv und unabhängig von dem Zeitpunkt oder speziellem Expertenwissen immer identisch.

Zur Anwendung des entwickelten Klassifikationsverfahrens im Feld muss das Verfahren im nächsten Schritt auf seine Robustheit gegenüber Reflexionen und Schattenbildungen getestet werden. D.h. es muss untersucht werden, ob das vorhandene 2 1/2 D-Modell zur Normalisierung der Farbinformationen und damit zur Eliminierung der Reflexionen und Schatten geeignet ist. Andernfalls müsste anhand eines anderen Verfahrens ein volles 3D-Modell erstellt werden. Eine erfolgreiche Anwendung des Verfahrens im Feld dürfte dann durch die Einsparung an Spritzmitteln deutliche ökologische und ökonomische Vorteile bieten.

Literaturverzeichnis

- Abraham, S. (1999). *Kamera Kalibrierung und metrische Auswertung monokularer Bildfolgen*. Ph. D. thesis, Universität Bonn, Institut für Geodäsie und Geoinformation, Bereich für Photogrammetrie.
- Albertz, J. (1999). *Einführung in die Fernerkundung*. Wissenschaftliche Buchgesellschaft.
- Asch, F. (2007). Eisentoxizität bei Reis – Merkmalsdefinition und Erkennung. *Vortr. Pflanzenzüchtg.* 72, S. 131–143.
- Bauer, S., F. Korč, und W. Förstner (2009). Investigation into the classification of diseases of sugar beet leaves using multispectral images. In E. van Henten, D. Goense, und C. Lokhorst (Hrsg.), *Precision Agriculture '09*, Wageningen, S. 229–238.
- Bauer, S., F. Korč, und W. Förstner (2011). The potential of automatic methods of classification to identify leaf diseases from multispectral images. *Precision Agriculture* 12, S. 361–377.
- Bell, G. E., D. L. Martin, S. G. Wiese, D. D. Dobson, M. W. Smith, M. L. Stone, und J. B. Solie (2002). Vehicle-mounted optical sensing: An objective means for evaluating turf quality. *Crop Science* 42, S. 197–201.
- Bilmes, J. A. (1998). A Gentle Tutorial of the EM Algorithm and its Application to Parameter Estimation for Gaussian Mixture and Hidden Markov Models. Technischer Bericht, University of California at Berkeley, International Computer Science Institute, Department of Electrical Engineering and Computer Science.
- Bishop, C. M. (2007). *Pattern Recognition and Machine Learning*. Springer.
- Chekuri, C., S. Khanna, J. Naor, und L. Zosin (2001). Approximation algorithms for the metric labelling problem via a new linear programming formulation. In *Symposium on Discrete Algorithms*.
- Cui, D., Q. Zhang, M. Li, G. Hartmann, und Y. Zhao (2010). Image processing methods for quantitatively detecting soybean rust from multispectral images. *Biosystems Engineering* 107, S. 186–193.

- Duda, R. O., P. E. Hart, und D. G. Stork (2001). *Pattern Classification*. Wiley-Interscience.
- Fink, G. A. (2003). *Mustererkennung mit Markov-Modellen: Theorie, Praxis, Anwendungsgebiete*. Teubner.
- Fukunaga, K. (1972). *Introduction to Statistical Pattern Recognition*. Academic Press.
- Gill, P. E., W. Murray, M. A. Saunders, und M. H. Wright (1984). Procedures for optimization problems with a mixture of bounds and general linear constraints. *ACM Trans. Math. Software* 10, S. 282 – 298.
- Gill, P. E., W. Murray, und M. H. Wright (1991). *Numerical Linear Algebra and Optimization*, Band 1. Addison Wesley.
- Hanan, J., B. Loch, und T. McAleer (2004). Processing laser scanner data to extract structural information. In C. Godin, J. Hanan, W. Kurth, A. Lacointe, A. Takenaka, P. Prusinkiewicz, T. DeJong, C. Beveridge, und B. Andrieu (Hrsg.), *Proceedings of the 4th International Workshop on Functional-Structural Plant Models*, Montpellier, S. 9–12.
- Hartley, R. und A. Zisserman (2003). *Multiple View Geometry in Computer Vision*. Cambridge University Press.
- Hillnhütter, C. und A.-K. Mahlein (2008). Neue Ansätze zur frühzeitigen Erkennung und Lokalisierung von Zuckerrübenkrankheiten. *Gesunde Pflanzen* 60, S. 143–149.
- Huang, K.-Y. (2007). Application of artificial neural network for detecting phalaenopsis seedling diseases using color and texture features. *Computers and electronics in agriculture* 57, S. 3–11.
- INPHO (2008). See the world with different eyes. Inpho: A Trimble Company.
- Jähne, B. (2005). *Digitale Bildverarbeitung*. Springer.
- Kang, S. B. und L. Quan (2010). *Image-Based Modeling of Plants and Trees*. Morgan & Claypool Publishers.
- Korč, F. und W. Förstner (2008). Approximate parameter learning in Conditional Random Fields: An empirical investigation. In G. Rigoll (Hrsg.), *Pattern Recognition*, Nummer 5096 in LNCS, S. 11–20. Springer.
- Korč, F. und D. Schneider (2007). Annotation Tool. Technischer Bericht, Dept. of Photogrammetry, University of Bonn.

- Koster, A. M. C. A., S. P. M. van Hoesel, und A. W. J. Kolen (1998). The partial constraint satisfaction problem: Facets and lifting theorems. *Operations Research Letters* 23(3-5), S. 89–97.
- Kraus, K. (2004). *Photogrammetrie: Geometrische Informationen aus Photographien und Laserscanneraufnahmen*. de Gruyter.
- Kumar, S., J. August, und M. Hebert (2005). Exploiting inference for approximate parameter learning in discriminative fields: An empirical study. In *5th International Workshop on Energy Minimization Methods in Computer Vision and Pattern Recognition*.
- Kumar, S. und M. Hebert (2006). Discriminative Random Fields. *International Journal of Computer Vision* 68(2), S. 179–201.
- KWS (2010). Blattkrankheiten: Schadbilder, Schäden, Strategien. KWS Saat AG.
- Läbe, T. und W. Förstner (2006). Automatic relative orientation of images. In *Proceedings of the 5th Turkish-German Joint Geodetic Days, March 29th - 31st*, Berlin.
- Lafferty, J., A. McCallum, und F. Pereira (2001). Conditional Random Fields: Probabilistic Models for Segmenting and Labeling Sequence Data. In *Proc. 18th International Conf. on Machine Learning*, S. 282–289.
- Lelong, C. C. D., P. C. Pinet, und H. Poilvé (1998). Hyperspectral imaging and stress mapping in agriculture: A case study on wheat in beauce. *Remote Sensing of Environment* 66 (2), S. 179–191.
- Lemaire, C. (2008). Aspects of the DSM production with high resolution images. In *Proceedings of the XXIst ISPRS congress*, Beijing, China.
- Li, S. Z. (2001). *Markov Random Field Modeling in Image Analysis*. Springer.
- Lowe, D. (2004). Distinctive image features from scale-invariant keypoints. *International Journal of Computer Vision* 60 (2), S. 91–110.
- Luhmann, T. (2003). *Nahbereichsphotogrammetrie: Grundlagen, Methoden und Anwendungen*. Wichmann.
- Ma, W. und H. Zha (2009). Convenient reconstruction of natural plants by images. In *19th International Conference on Pattern Recognition*, S. 1–4.
- Mahlein, A.-K., U. Steiner, H.-W. Dehne, und E.-C. Oerke (2010). Spectral signatures of sugar beet leaves for the detection and differentiation of diseases. *Precision Agriculture* 11, S. 413–431.

- McGlone, J. C., E. M. Mikhail, J. Bethel, und R. Mullen (Hrsg.) (2004). *Manual of Photogrammetry*. American Society for Photogrammetry and Remote Sensing.
- Meier, U., L. Bachmann, H. Buhtz, H. Hack, R. Klose, und M. B. et al. (1993). Phänologische Entwicklungsstadien der Beta-Rüben (*Beta vulgaris* L. ssp.). Codierung und Beschreibung nach der erweiterten BBCH-Skala (mit Abbildungen). *Nachrichtenblatt Deutscher Pflanzenschutzdienst* 45, S. 37–41.
- Mewes, T., B. Waske, J. Franke, und G. Menz (2010). Derivation of stress severities in wheat from hyperspectral data using support vector regression. In *Proc. IEEE-Whispers*, Reykjavik, Island, S. 1–4.
- Oommen, B. J., J. R. Zgierski, und D. Nussbaum (2004). Deterministic majority filters applied to stochastic sorting. In *Proceedings of the 42nd annual Southeast regional conference*.
- Pan, Z., W. Hu, X. Guo, und C. Zhao (2004). An efficient image-based 3D reconstruction algorithm for plants. In T. Kanade, J. Kittler, J. Kleinberg, F. Mattern, J. Mitchell, O. Nierstrasz, C. Pandu Rangan, B. Steffen, M. Sudan, D. Terzopoulos, D. Tygar, M. Vardi, und G. Weikum (Hrsg.), *Lecture Notes in Computer Science*, Heidelberg, S. 751–760. Springer.
- Poynton, C. (2003). *Digital Video and HDTV Algorithms and Interfaces*. Morgan Kaufmann Publishers.
- Pydipati, R., T. F. Burks, und W. S. Lee (2006). Identification of citrus disease using color texture features and discriminant analysis. *Computers and electronics in agriculture* 52, S. 49–59.
- Rösch, C., M. Dusseldorp, und R. Meyer (2006). *Precision Agriculture: Landwirtschaft mit Satellit und Sensor*. Deutscher Fachverlag.
- Sanyal, P. und S. C. Patel (2008). Pattern recognition method to detect two diseases in rice plants. *The Imaging Science Journal* 56, S. 319–325.
- Schlesinger, M. I. (1976). Sintaksicheskiy analiz dvumernykh zritelnykh signalov v usloviyakh pomekh (Syntactic analysis of two-dimensional visual signals in noisy conditions). *Kibernetika* 4, S. 113–130.
- Steddom, K., M. W. Bredehoeft, M. Khan, und C. M. Rush (2005). Comparison of visual and multispectral radiometric disease evaluations of cercospora leaf spot of sugar beet. *Plant Disease* 89, S. 153–158.
- Stuppy, W., J. Maisano, M. Colbert, P. Rudall, und T. Rowe (2003). Three-dimensional analysis of plant structure using high-resolution X-ray computed tomography. *Trends in Plant Science* 8, S. 2–6.

- The MathWorks, I. (Hrsg.) (2010). *Optimization Toolbox 5 - User's Guide*. MathWorks.
- Vincent, L. und P. Soille (1991). Watersheds in digital spaces: An ancient algorithm based on immersion simulations. *IEEE Transactions on Pattern Analysis and Machine Intelligence* 13(6), S. 583–598.
- Wainwright, M. J., T. S. Jaakkola, und A. S. Willsky (2005). Map estimation via agreement on trees: message-passing and linear programming. *IEEE Transactions on Information Theory* 51(11), S. 3697–3717.
- Wolf, P. F. J. und J.-A. Verreet (1997). Epidemiologische Entwicklung von *Cercospora beticola* (Sacch.) in Zuckerrüben. *Journal of Plant Diseases and Protection* 104 (6), S. 545–556.
- Wolf, P. F. J. und J.-A. Verreet (2002). The IPM Sugar Beet Model: An Integrated Pest Management System in Germany for the Control of Fungal Leaf Diseases in Sugar Beet. *Plant Disease* 86 (4), S. 336–344.