

GeoBerichte 54

LANDESAMT FÜR
BERGBAU, ENERGIE UND GEOLOGIE



Erkennung von
Entwässerungsgräben
mit Methoden des
Maschinellen Lernens



Niedersachsen



GeoBerichte 54

Landesamt für
Bergbau, Energie und Geologie

Erkennung von Entwässerungsgräben mit Methoden des Maschinellen Lernens

JOST WESSELS, MITHRA-CHRISTIN HAJATI &
JÖRG ELBRACHT

Hannover 2025

Impressum

Herausgeber: © Landesamt für Bergbau, Energie und Geologie

Stilleweg 2
30655 Hannover
Tel. (0511) 643-0

Download unter www.lbeg.niedersachsen.de

1. Auflage

Version: 01.10.2025

Redaktion: Ricarda Nettelmann

Mail: bodenkundlicheberatung@lbeg.niedersachsen.de

Titelbild: Fotos: J. Wessels (links), M. Hoetmer (Mitte), G. Griffel (rechts).

ISSN 1864–6891 (Print)

ISSN 1864–7529 (digital)

DOI 10.48476/geober_54_2025

GeoBer.	54	S. 3 – 51	29 Abb.	6 Tab.	Anh.	Hannover 2025
---------	-----------	-----------	---------	--------	------	---------------

Erkennung von Entwässerungsgräben mit Methoden des Maschinellen Lernens

JOST WESSELS, MITHRA-CHRISTIN HAJATI & JÖRG ELBRACHT

Kurzfassung

In diesem Geobericht wird eine Methode zur Erkennung von Entwässerungsgrabenstrukturen auf Grundlage des Digitalen Geländemodells (DGM1) für Niedersachsen vorgestellt. Bei der Modellierung kamen unter anderem zwei Methoden aus dem Bereich des Maschinellen Lernens (ML) zum Einsatz.

Zunächst wurden auf Basis von zuvor digitalisierten Gräben sowie von mehreren aus dem DGM1 abgeleiteten Terrainindices Random-Forest-Modelle in mehreren Pilotgebieten trainiert. Nach einer Evaluation und Anwendung auf die jeweils anderen Pilotgebiete wurde das am besten geeignete Modell ausgewählt und auf die gesamte Fläche des Landes Niedersachsen angewendet. Im Rahmen eines anschließenden Post-Processings wurden die Zwischenergebnisse mit weiteren Daten angereichert und anhand zweier XGBoost-Modelle gefiltert. Die Ergebnisse können zum einen Hinweise auf noch nicht erfasste Entwässerungsgräben liefern. Zum anderen besteht die Möglichkeit, bereits bekannte Strukturen hinsichtlich ihrer tatsächlichen Lage im Gelände mit den modellierten Ergebnissen abzugleichen. Der vorliegende Bericht folgt im Wesentlichen den einzelnen Arbeitsschritten innerhalb des entwickelten Workflows, um eine möglichst gute Nachvollziehbarkeit zu gewährleisten.

Die Entwicklung und Durchführung der Methode erfolgte im Rahmen des LBEG-Projektes „Bodenwasserhaushalt und Grundwassermenge im Klimawandel“ (Teilprojekt KliBoG1), welches durch das Niedersächsische Kompetenzzentrum Klimawandel (NIKO) des Niedersächsischen Ministeriums für Umwelt, Energie, Bauen und Klimaschutz (MU) gefördert wird. Innerhalb des Projektes sollen die Modellergebnisse unter anderem dazu genutzt werden, die Datengrundlage in Bezug auf künstlich entwässerte Flächen zu verbessern. Hierzu sind weitere Analysen und Modellierungen vorgesehen. Die Ergebnisse der vorliegenden Untersuchungen sind im *NIBIS®-Kartenserver* veröffentlicht.

Inhalt

Vorwort.....	5
1. Hintergrund.....	6
2. Datengrundlage und Methodik	6
2.1. Übersicht zur Grabenerkennung.....	6
2.2. Datengrundlage und Modellansätze	7
3. Vom DGM zur Grabenwahrscheinlichkeit.....	10
3.1. Auswahl der Trainingsgebiete sowie Digitalisierung von Gräben.....	10
3.2. Auswahl der Prädiktorvariablen	12
3.3. Pre-Processing der Trainingsdatensätze.....	15
3.4. Aufbau von Random-Forest-Modellen und Auswahl des besten Modells.....	16
3.5. Parametrisierung und Interpretation des finalen Random-Forest-Modells.....	17
3.6. Zwischenergebnis nach Anwendung des finalen Random-Forest-Modells.....	18
4. Clustern von Grabenstrukturen.....	21
4.1. Ausschlusskriterien	21
4.2. Entrauschen und Clustern	22
5. Filtern der Graben-Cluster	23
5.1. Anreichern der Cluster mit weiteren Informationen	24
5.2. Labeling von Clustern und Entwicklung von XGBoost-Modellen.....	25
5.3. Parametrisierung und Modellinterpretation	26
6. Ergebnisse nach Durchführung des gesamten Workflows.....	32
6.1. Ergebnis nach Anwendung der XGBoost-Modelle	32
6.2. Abschätzung der Modellgüte für den gesamten Workflow	32
6.3. Verteilung der modellierten Strukturen in der Fläche	36
7. Diskussion und Ausblick	40
7.1. Zusammenfassung.....	40
7.2. Möglichkeiten und Grenzen	42
7.3. Nutzungshinweise und FAQ	44
8. Dank.....	45
9. Quellen	46
Anhang.....	50

Vorwort

Liebe Leserinnen und Leser,

Wasser ist die Grundlage allen Lebens – eine Wahrheit, die aktueller denn je erscheint. Für uns Menschen ist Wasser nicht nur das wichtigste Lebensmittel, sondern auch unverzichtbar für Landwirtschaft, Industrie und die Funktionsfähigkeit unserer Ökosysteme. Verfügbarkeit und Qualität von Wasser bestimmen maßgeblich unsere Lebensbedingungen und sind eng mit Gesundheit, Wohlstand und ökologischer Stabilität verknüpft.

Eine herausragende Stellung nimmt hierbei das Grundwasser ein. In Niedersachsen stellt es in weiten Teilen die wichtigste Ressource für die Trinkwasserversorgung dar und sichert so die Versorgung von Millionen von Menschen. Doch Grundwasser ist keine unerschöpfliche Quelle. Seine Neubildung – also die natürliche Auffüllung der unterirdischen Wasserspeicher durch versickerndes Niederschlagswasser – ist ein empfindlicher Prozess, der von zahlreichen natürlichen und anthropogenen Faktoren beeinflusst wird.

Ein entscheidender Aspekt ist hierbei die Entwässerung unserer Landschaft. Wo Wasser durch technische Maßnahmen wie Gräben oder Drainagen gezielt abgeleitet wird, verringert sich die Möglichkeit zur Grundwasserneubildung. Bei hoch anstehendem Grundwasser kann es sogar zu einer dauerhaften Grundwasserzehrung kommen. Die Auswirkungen solcher Entwässerungsmaßnahmen sind vielfältig und betreffen nicht nur die Wasserversorgung, sondern auch den Naturhaushalt, den Boden- und Gewässerschutz sowie die landwirtschaftlichen Rahmenbedingungen.

Trotz der großen Bedeutung dieses Themas liegen für Niedersachsen bislang keine flächendeckenden Kenntnisse über das tatsächliche Ausmaß und die Verteilung von Drainagen oder Entwässerungsgräben vor. Mit der vorliegenden Arbeit wird ein wichtiger erster Schritt unternommen, um diese Wissenslücke zu schließen. Durch den Einsatz moderner Methoden des maschinellen Lernens wird die automatisierte Erkennung von Entwässerungsgräben ermöglicht und damit eine belastbare Datengrundlage geschaffen.

Diese bildet künftig die Basis für weitere Untersuchungen und Modellierungen – unter anderem zur verbesserten Berechnung der Grundwasserneubildung mit dem am LBEG eingesetzten Modell mGROWA. Darüber hinaus eröffnen sich weitere Anwendungsbereiche, etwa beim Wasserrückhalt in der Fläche, beim Hochwasserschutz sowie bei der Aufnahme von durch Erosion verlagerten Bodenpartikeln. Damit leistet diese Arbeit einen wichtigen Beitrag zum besseren Verständnis der Landschaftsentwässerung – einem Thema, das vor dem Hintergrund des Klimawandels und der erforderlichen Anpassungsmaßnahmen zunehmend an Bedeutung gewinnt.

Ich danke allen Beteiligten für ihre engagierte Arbeit und wünsche den Leserinnen und Lesern eine aufschlussreiche Lektüre.

Carsten Mühlenmeier

Präsident LBEG



1. Hintergrund

Die künstliche Entwässerung in Niedersachsen, insbesondere von landwirtschaftlich genutzten Flächen, stellt einen wichtigen Einflussfaktor für den Boden- und Landschaftswasserhaushalt dar (ERTL et al. 2019, GEHRT et al. 2019). Neben den positiven Auswirkungen wie beispielsweise einer besseren Befahrbarkeit im Frühjahr und einer verbesserten Ertragsfähigkeit ansonsten vernässter Böden können Drainagesysteme auch zu veränderten Spitzenabflüssen bzw. Hochwassergefahren führen (GRAMLICH et al. 2018). Des Weiteren wird die Grundwasserneubildung durch künstliche Entwässerung reduziert. Durch die Anlage von Drainageleitungen und/oder Entwässerungsgräben wird Sickerwasser zum Teil bereits vor dem Erreichen des Grundwassers in die Vorflut abgeführt und steht somit nicht mehr für die Bildung von neuem Grundwasser zur Verfügung (GEHRT et al. 2019). Zudem kann auch bereits vorhandenes Grundwasser durch künstliche Entwässerungsmaßnahmen drainiert werden. Dies ist in Niedersachsen vor allem in Niederungsbereichen der Fall (GEHRT et al. 2019, TETZLAFF et al. 2008). In der Bilanz kann somit teilweise nicht nur im Sommer, sondern ganzjährig eine negative Grundwasserneubildung, d. h. Grundwasserzehrung stattfinden (ERTL et al. 2024).

Die potenzielle Dränfläche zur Regulierung des Grundwasserstandes wird von GEHRT et al. (2019) beispielsweise für das Küstenholozän mit 59 % und für die Flusslandschaften mit 53 % angegeben. Auch für die Geestbereiche sowie die Flachlandschaften des Bergvorlandes werden mit 38 % bzw. 22 % noch vergleichsweise hohe Anteile potenziell dränkter Flächen ausgewiesen (GEHRT et al. 2019). Insbesondere im Bergvorland kommt ergänzend ein hoher Anteil an Flächen, welche eine sogenannte Bedarfsdränage zur Regulation von Staunässe aufweisen, hinzu (GEHRT et al. 2019).

Zwar wird die Anlage von Drainagesystemen beispielsweise in der Landwirtschaft in der Regel dokumentiert, etwa im Zuge einer Flurbereinigung. Auch werden bestehende Entwässerungsgräben bzw. Gewässer dritter Ordnung zumeist in regelmäßigen oder unregelmäßigen Abständen unterhalten und gepflegt und sind daher häufig nicht unbekannt. Eine systematische Erfassung von Entwässerungsstrukturen gibt es in Niedersachsen jedoch nicht. Im Pro-

jekt „Bodenwasserhaushalt und Grundwassermenge im Klimawandel“ (Teilprojekt KliBoG1) soll daher unter anderem die Datenbasis zu künstlich entwässerten Flächen für die Modellierung der Grundwasserneubildung verbessert werden. In einem ersten Schritt werden vorhandene Entwässerungsgrabenstrukturen, welche bislang nur teilweise kartiert bzw. digitalisiert wurden, aus dem Digitalen Geländemodell abgeleitet. Auf die Methodik sowie die Ergebnisse der vorgenannten Modellierung wird nachfolgend näher eingegangen. Die Ergebnisse sollen unter anderem innerhalb des KliBoG1-Teilprojektes für weitergehende Analysen herangezogen werden, um schließlich die maßgebliche Hintergrundinformation zur künstlichen Entwässerung für die landesweite Wasserhaushaltsmodellierung mit mGROWA (ERTL et al. 2024) zu verbessern.

2. Datengrundlage und Methodik

2.1. Übersicht zur Grabenerkennung

Für die Ableitung natürlicher Gewässerverläufe und -netze aus rasterbasierten Digitalen Geländemodellen (DGM) werden traditionellerweise Algorithmen zur Ermittlung von Fließrichtung (flow direction) und Abflussakkumulation (flow accumulation) verwendet (STANISLAWSKI et al. 2018). Zur Erkennung von künstlichen Entwässerungsstrukturen in anthropogen überprägten Landschaften sind die vorgenannten Ansätze hingegen zumeist nicht zielführend. Zum einen orientieren sich die Verläufe von Entwässerungsgräben neben der natürlichen großräumigen Topographie vor allem an anderen räumlichen Einheiten wie Straßen oder landwirtschaftlichen Flächen (BAILLY et al. 2008, BUCHANAN et al. 2013, CAZORZI et al. 2013). Zum anderen werden insbesondere Niederungsbereiche mit geringer Reliefenergie künstlich entwässert, in denen die hydrologischen Ansätze häufig zu ungenauen Ergebnissen führen (CAZORZI et al. 2013, DU et al. 2024).

Daher kommen bei der automatisierten Erkennung von Entwässerungsgräben verstärkt lokale topographische Parameter als Eingangsvariablen zum Einsatz. Teilweise werden hierfür mittels Laserscanning (Light Detection and

Ranging, LiDAR) erzeugte Punktwolken verwendet (BAILLY et al. 2008; BALADO et al. 2019; ROELEN, HÖFLE et al. 2018). Überwiegend werden jedoch die aus diesen Punktwolken abgeleiteten DGM als Eingangsdaten genutzt (CAZORZI et al. 2013; FLYCKT et al. 2022; GRAVES et al. 2020; LIDBERG et al. 2023; PASSALACQUA et al. 2012; QIAN et al. 2018; RAPINEL et al. 2015; ROELEN, ROSIER et al. 2018). Bei entsprechender Datenlage können alternativ oder ergänzend auch aus optischen Daten abgeleitete Eingangsvariablen in die Modellierung eingehen (DU et al. 2024; ROBB et al. 2023; ROELEN, HÖFLE et al. 2018). Auch die Einbeziehung der beim Laserscanning ermittelten Reflexionsintensität ist prinzipiell möglich (BROERSEN et al. 2017; DU et al. 2024; ROELEN, HÖFLE et al. 2018).

Methodisch kommen im Zuge der Modellierung von Gewässer- bzw. Grabenstrukturen neben der Ableitung anhand von Schwellenwerten (GRAVES et al. 2020, PASSALACQUA et al. 2012, QIAN et al. 2018, RAPINEL et al. 2015) zunehmend Verfahren des Maschinellen Lernens (ML) zum Einsatz (BALADO et al. 2019; DU et al. 2024; FLYCKT et al. 2022; LIDBERG et al. 2023; ROBB et al. 2023; ROELEN, ROSIER et al. 2018). Die Güte der Ergebnisse hängt hierbei neben dem gewählten Modellansatz vor allem von der Qualität der Eingangsparameter sowie von deren Auswahl ab. Des Weiteren ist die Auswahl

der Trainings- und Testgebiete für die Modellgüte von großer Bedeutung. Hierbei steigen die Herausforderungen und letztlich auch die Fehleranfälligkeit bzw. die Ungenauigkeit der Prognose tendenziell mit der Komplexität der Topographie sowie mit einem zunehmenden Baumbestand (vgl. CAZORZI et al. 2013, FLYCKT et al. 2022).

2.2. Datengrundlage und Modellansätze

Es wurde zunächst davon ausgegangen, dass Entwässerungsgräben bestimmte morphologische Eigenschaften aufweisen, durch welche sie von anderen Strukturen in der Landschaft unterschieden werden können. So liegt beispielsweise die Grabensohle tiefer als ihre unmittelbare Umgebung (Abb. 1a). Des Weiteren werden Gräben auf zwei gegenüberliegenden Seiten von Böschungen begrenzt, wobei das Querprofil idealerweise V-förmig ist (Abb. 1b). Die Grabenböschungen führen außerdem dazu, dass an der Grabensohle eine Horizontbegrenzung vorliegt und die „Sicht“ somit mehr oder weniger stark eingeschränkt ist (Abb. 1c). In der Draufsicht stellen Entwässerungsgräben linienhafte Elemente mit einem zumeist geraden Verlauf dar, welche Wasser von seitlich angrenzenden Flächen aufnehmen (Abb. 1d).

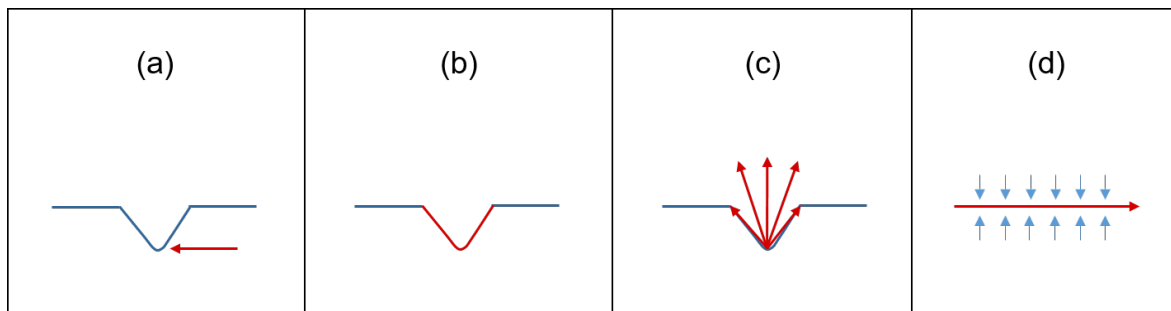


Abb. 1: Morphologische Eigenschaften von Entwässerungsgräben zur Unterscheidung von anderen Elementen in der Landschaft (a bis c: Querschnitt, d: Draufsicht). Erläuterungen s. Text.

Um die vorgenannten Eigenschaften möglichst gut abbilden zu können, wurden verschiedene Terrainindices identifiziert, welche sich auf Basis des Digitalen Geländemodells ermitteln lassen. Bei einem Digitalen Geländemodell (DGM) handelt es sich um einen Rasterdatensatz, bei dem jeder Rasterzelle ein Höhenwert zugeordnet ist. Im vorliegenden Fall wurde ein DGM mit einer Rasterweite von 1 m verwendet (DGM1). Auf die verschiedenen Indices wird im Abschnitt 3.2 genauer eingegangen. Der Modellansatz basiert auf der Annahme, dass eine Rasterzelle, die zu einem Graben gehört, andere charakteristische Werte bzw. Wertespektren der jeweiligen Terrainindices aufweist, als eine Rasterzelle, welche beispielsweise einer ebenen Fläche oder einem Kuppenbereich zuzuordnen ist. Im vorliegenden Fall wurde hierzu ein Modell auf Basis des Random-Forest-Algorithmus‘ (BREIMAN 2001) entwickelt.

Die Zwischenergebnisse dieses Random-Forest-Modells wurden anschließend einem mehrstufigen Post-Processing unterzogen. Hierbei

wurden unter anderem mehrere Ausschlusskriterien für das Auftreten von Entwässerungsgräben definiert. Als wichtigster Bestandteil wurde schließlich ein auf dem XGBoost-Algorithmus (CHEN & GUESTRIN 2016) basierendes Modell entwickelt, mit welchem vor allem nichtlineare Strukturen aus den Zwischenergebnissen herausgefiltert wurden. Als Prädiktorvariablen kamen hierbei neben der Linearität (bzw. Kompaktheit) der Strukturen auch weitere Parameter zum Einsatz, welche zum Teil aus weiteren Datenquellen extrahiert wurden. Hierauf wird im Abschnitt 5.1 noch näher eingegangen. Die Tabelle 1 zeigt eine Übersicht aller für die vorliegende Arbeit genutzten Datensätze inklusive der jeweiligen Quellen und der Verwendung.

Die Abbildung 2 zeigt einen schematischen Ablauf des im vorliegenden Bericht angewendeten Workflows. In den nachfolgenden Abschnitten wird auf die einzelnen Schritte genauer eingegangen. Die Berechnungen erfolgten nahezu durchgehend mit dem Statistikpaket R (R CORE TEAM 2023) bzw. RStudio (POSIT TEAM 2023).

Tab. 1: Verwendete Datensätze, ihre Quellen und Verwendung.

Datensatz	Quelle	Verwendung
DGM1	LGLN 2024b	- Überprüfung bei der Digitalisierung von Entwässerungsgräben und Plausibilitätsprüfung der Ergebnisse, - Berechnung von Terrainindices, - Extraktion der Geländehöhe, - Berechnung von Geländeformen.
Linienhafte Gewässer	LGLN 2024c	Überprüfung bei der Digitalisierung von Entwässerungsgräben und Plausibilitätsprüfung der Ergebnisse.
Digitale Orthophotos (DOP)	LGLN 2024a	Überprüfung bei der Digitalisierung von Entwässerungsgräben und Plausibilitätsprüfung der Ergebnisse.
Feldblöcke	SLA 2022	Definition von Ausschlusskriterien.
Flächenhafte Gewässer	LGLN 2024c	Definition von Ausschlusskriterien.
Bodenkarte 1 : 50.000	GEHRT et al. 2021	- Definition von Ausschlusskriterien, - Extraktion von Bodenregionen.
Waldflächen	LGLN 2024c	Extraktion von bewaldeten Gebieten.

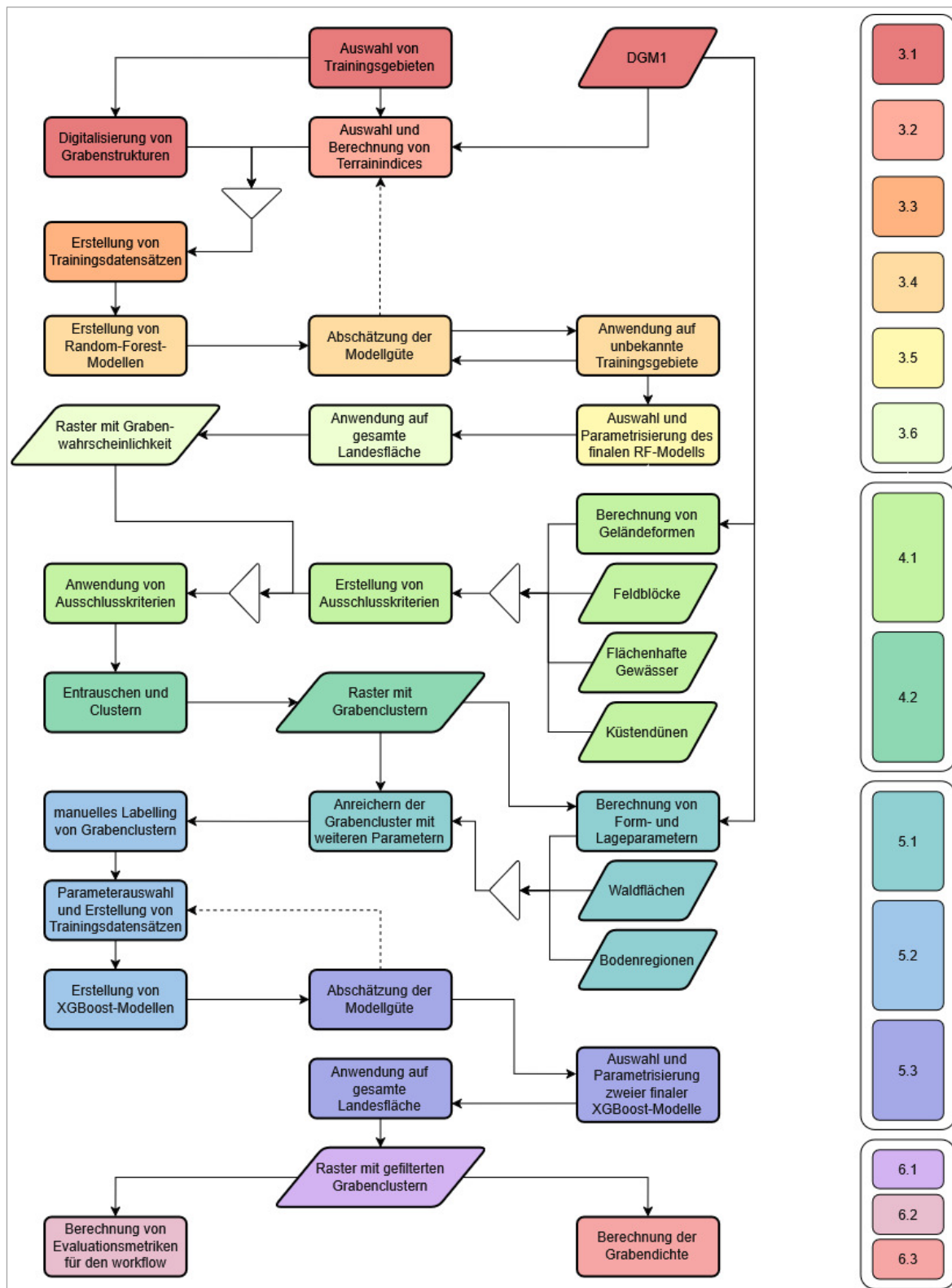


Abb. 2: Schematischer Ablauf der Modellierung. Auf der rechten Seite sind die zugehörigen Textabschnitte angegeben.

3. Vom DGM zur Grabenwahrscheinlichkeit

In diesem Abschnitt wird ein Random-Forest-Modell (BREIMAN 2001) beschrieben, welches anhand von verschiedenen Terrainindices sowie bekannten Grabenverläufen trainiert wird. Das Ergebnis nach Anwendung des Modells ist ein Raster mit einer Zellengröße von 1 x 1 m, in dem jeder Zelle in Niedersachsen eine „Grabenwahrscheinlichkeit“ zugewiesen wird.

3.1. Auswahl der Trainingsgebiete sowie Digitalisierung von Gräben

Zunächst wurden insgesamt vier Trainingsgebiete zu je 4 km² innerhalb des Einzugsgebietes des Dümmers ausgewählt. Zum einen sind die drainierten Flächen und die vorhandenen Entwässerungsgräben in diesem Bereich durch regelmäßige Probenahmen und Geländebegehungen im Rahmen von anderen Projekten des LBEG vergleichsweise gut bekannt (RÖDER 2022, RÖDER & SCHÄFER 2015). Zum anderen

ist das Gebiet mit den ausgedehnten Talsandniederungen, Moorflächen und Geesten sowie – im Süden des Einzugsgebietes – den Lössböden und schließlich den Höhenzügen des Wiehengebirges durch verschiedene topographische Einheiten gekennzeichnet. In diesen Trainingsgebieten wurden die tatsächlich vorhandenen Gräben digitalisiert, d. h. in ihrem Verlauf in einem GIS manuell erfasst. Dies erfolgte unter Berücksichtigung des DGM1 (Schummerung), des vorhandenen digitalen Gewässernetzes sowie anhand von Luftbildern bzw. Digitalen Orthophotos. Im Nachgang wurden weitere Gebiete ausgewählt, um weitere für Niedersachsen typische Landschaftsräume mit einzubeziehen. Fünf dieser zusätzlichen Gebiete wurden schließlich ebenfalls für das Modelltraining berücksichtigt. Die Abbildung 3 zeigt die Lage dieser insgesamt neun Trainingsgebiete in Niedersachsen sowie einen Detailausschnitt der vier Trainingsgebiete im Einzugsgebiet des Dümmers. Ein zwischenzeitlich in Betracht gezogener Trainingsgebiet mit der Nummer 9 wurde nicht verwendet.

In der Tabelle 2 sind einige Eigenschaften der Trainingsgebiete zusammengestellt.

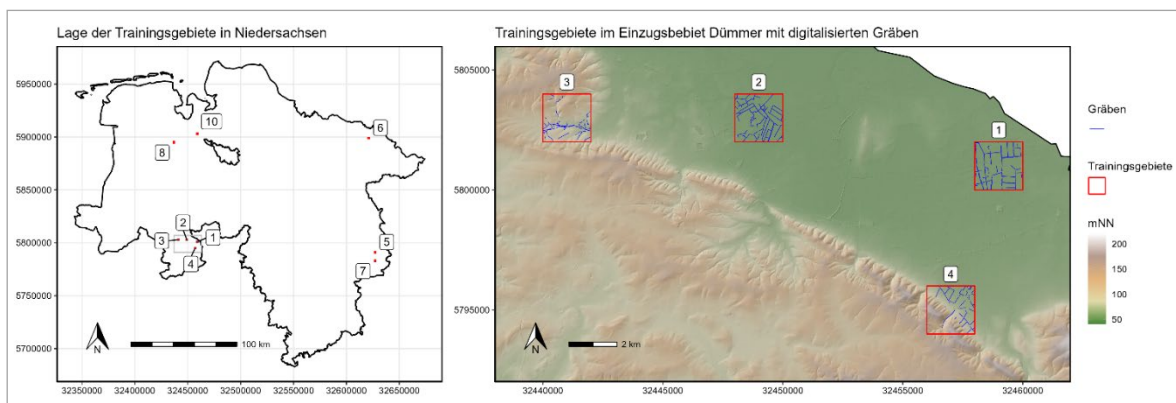


Abb. 3: Übersicht zur Lage der insgesamt neun Trainingsgebiete in Niedersachsen (links) sowie Ausschnitt aus dem Einzugsgebiet Dümmers mit den Trainingsgebieten 1 bis 4 inklusive digitalisierter Gräben (rechts).

Tab. 2: Ausgewählte Eigenschaften der neun Trainingsgebiete.

Nr.	Bezeichnung	Koordinaten*	Höhenlage [NHN]	Bodenregion**	Hauptsächliche Nutzung
1	Wimmerbach	32458000, 5800000	43,7 – 48,9	Geest	Landwirtschaft
2	Kronensee	32448000, 5802000	42,8 – 50,6	Geest	Landwirtschaft, Wald
3	Venner Mühlenbach	32440000, 5802000	67,0 – 143,2	Geest, Bergland, Bergvorland	Landwirtschaft, Wald
4	Hüsedede	32456000, 5794000	51,2 – 197,5	Bergvorland, Bergland	Landwirtschaft, Wald, Siedlung
5	Schickelsheim	32626000, 5790000	94,9 – 123,2	Bergland, Geest, Flusslandschaften, Bergvorland	Landwirtschaft
6	Schieringen	32620000, 5898000	51,6 – 107,5	Bergland	Wald, Landwirtschaft
7	Raebke	32626000, 5782000	164,3 – 229,5	Geest	Wald, Landwirtschaft
8	Aschhauserfeld	32436000, 5894000	5,3 – 14,5	Geest	Landwirtschaft, Wald
10	Alte Kapelle	32458000, 5902000	-1,6 – 1,8	Küstenholozän	Landwirtschaft
* x- und y-Wert der südwestlichen Ecke (UTM32, epsg 4647), ** gemäß Bodenkarte 1 : 50.000 (GEHRT et al. 2021).					

Bei der Digitalisierung der Grabenverläufe wurde insbesondere darauf geachtet, die Grabensohle zu erfassen. Vorrangig wurden landwirtschaftliche Entwässerungsgräben digitalisiert (s. Abbildung 4, Mitte sowie rechts). Aufgrund einer ähnlichen Morphologie wurden auch Gräben berücksichtigt, die beispielsweise

der Entwässerung von Verkehrs- oder Waldflächen dienen oder die mehrere Funktionen erfüllen (s. Abbildung 4, links).

In der Abbildung 5 ist exemplarisch ein kleiner Ausschnitt aus dem Trainingsgebiet 1 (Wimmerbach) mit dem Verlauf der digitalisierten Gräben dargestellt.

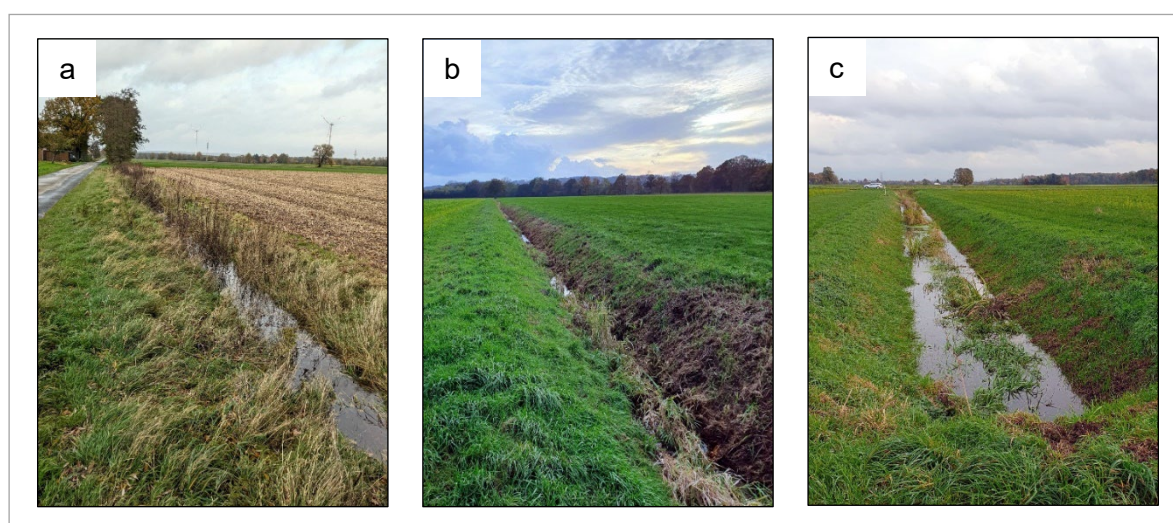


Abb. 4: Beispiele für Grabenstrukturen, die für das Modelltraining digitalisiert wurden. a: Trainingsgebiet 1 (Wimmerbach), b und c: Trainingsgebiet 2 (Kronensee). Fotos: C. Röder (a), M. Hoetmer (b, c).

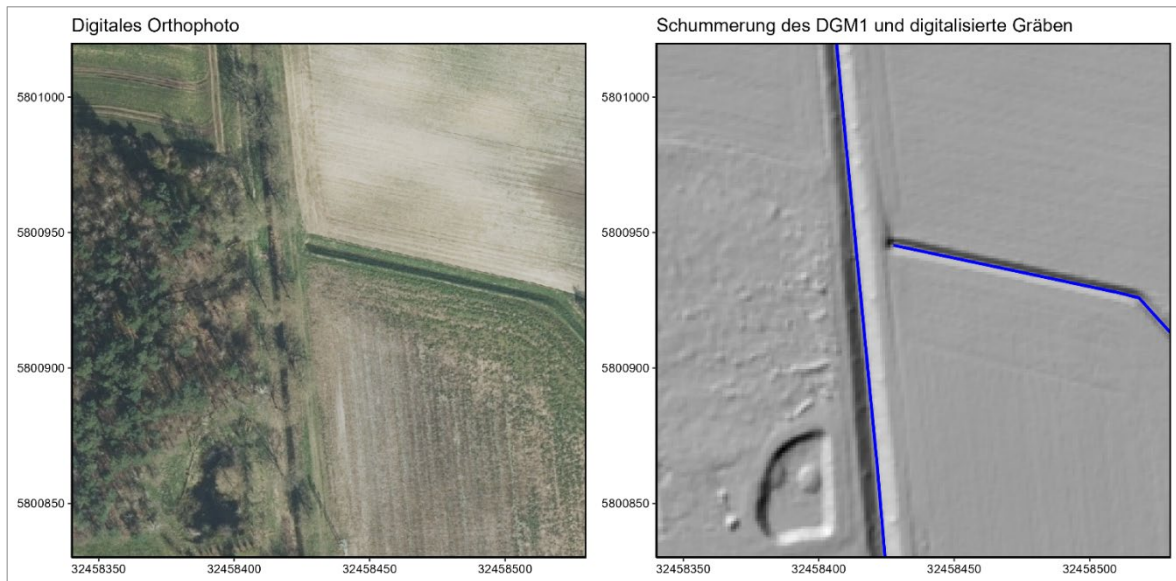


Abb. 5: Ausschnitt aus dem Trainingsgebiet 1 (Wimmerbach) mit dem zugehörigen Digitalen Orthophoto (links) und der DGM1-Schummerung inklusive der manuell digitalisierten Grabenverläufe (rechts).

3.2. Auswahl der Prädiktorvariablen

Im Zuge der Modellierung wurden verschiedene Terrainindices als mögliche Prädiktorvariablen getestet. Insbesondere die im Abschnitt 2.2 bzw. in der Abbildung 1 genannten Eigenschaften von Entwässerungsgräben sollten hierbei abgebildet werden. Es wurde jeweils geprüft, ob eine Verbesserung der Modellergebnisse durch das Einbeziehen der jeweiligen Variable erreicht werden konnte. Sofern eine Variable zu keiner nennenswerten Verbesserung beitrug, wurde sie verworfen. Hierzu gehören unter anderem die Hangneigung sowie der Terrainindex „Curvature“. Bei relevanten Variablen bzw. Indices wurden unterschiedliche Parametrisierungen getestet. Je Index wurden schließlich zwei bzw. drei unterschiedliche „Varianten“ ausgewählt. Nachfolgend werden diese tatsächlich verwendeten Indices näher erläutert.

- **Difference from mean Elevation (DME):**
Dieser Index berechnet die relative Höhenlage einer Rasterzelle im Vergleich zu ihrer unmittelbaren Umgebung (LINDSAY 2016). Ein positiver Wert zeigt an, dass die Zelle höher liegt, als der Mittelwert ihrer Nachbarzellen; ein negativer Wert hingegen gibt an, dass die Zelle tiefer als ihre lokale Umgebung liegt. Rasterzellen, die zu einer Grabensohle gehören, weisen daher – je nach Größe des Suchradius und der Breite der Sohle – in der Regel negative Werte auf.
- **Laplace of Gauss Filter (LGF):**
Der Laplace-Filter ermittelt die zweite Ableitung der Höhe aus dem DGM und stellt somit ein Maß für die Änderung der Hangneigung dar. Um die Anfälligkeit gegenüber kleinräumiger Variabilität der Geländeoberfläche zu reduzieren, wird der Berechnung eine Glättung („smoothing“) des DGM mit Hilfe eines Gauss-Filters vorangestellt (KOKALJ & HESSE 2017). In einer Ebene liegt der Index dementsprechend bei 0 (LINDSAY 2016). Im vorliegenden Fall werden im Allgemeinen für die Bereiche der Böschungsoberkanten positive Werte ausgewiesen, während die Grabensohlen negative Werte aufweisen.

■ **Impoundment Size Index (ISI):**

Bei diesem Index wird für jede Rasterzelle ein potenzielles Staubecken berechnet, welches sich ergäbe, wenn an dieser Rasterzelle ein Staudamm mit einer bestimmten Länge („dam length“) errichtet würde (LINDSAY 2016, WU & BROWN 2022). Hierbei ist von Vorteil, dass dies nur bei vorhandenen seitlichen Begrenzungen wie etwa Böschungen funktioniert und somit ein weiteres im Abschnitt 2.2 genanntes Merkmal von Entwässerungsgräben abgebildet werden kann. Gleichzeitig können breitere Geländeeinschnitte durch die Wahl einer geeigneten Dammlänge („dam length“) ausgeschlossen werden. Bei der Berechnung des Impoundment Size Index werden je Zelle zunächst vier Dämme mit unterschiedlicher Ausrichtung (Nord-Süd, Nordost-Südwest, Ost-West, Südost-Nordwest) simuliert (GUDIM 2019). Der Damm, welcher die größte Kronenhöhe aufweist, wird schließlich für die Ergebnisausgabe verwendet (GUDIM 2019). Neben der Größe des Staubeckens kann auch die Höhe des potenziellen Staudamms („dam height“) als Ergebnis ausgegeben werden. Dies wurde im vorliegenden Fall durchgeführt, sodass mit dem Impoundment Size Index ein Maß für die Tiefe eines etwaigen Geländeeinschnittes an jeder Rasterzelle vorliegt. Die Ausgabewerte der „dam height“ fallen stets ≥ 0 aus und steigen, je stärker die Vertiefung im Gelände ist.

■ **Sky View Factor (SVF):**

Dieser Index repräsentiert den Anteil des Himmels, welcher von einem bestimmten Punkt aus sichtbar ist und nimmt daher Werte zwischen 0 (kein Himmel sichtbar) und 1 (Himmel vollständig sichtbar) an (KOKALJ et al. 2011). Bei der Anwendung auf ein Digitales Geländemodell erhalten beispielsweise Rasterzellen, welche auf einem Gipfel liegen, einen Wert von 1, während für Rasterzellen in einer tiefen Schlucht Werte nahe 0 ausgewiesen werden. Bei der Berechnung des Sky View Factors wird an dem jeweiligen Punkt eine oberhalb liegende virtuelle Halbkugel betrachtet (KOKALJ & HESSE 2017). Durch die Wahl eines Suchradius kann definiert werden, bis zu welcher Entfernung etwaige Horizontbegrenzungen berücksichtigt werden sollen. Im vorliegenden Fall werden durch vergleichsweise geringe Radien die Böschungen eines Grabens einbezogen, während weiter entfernte Hindernisse (Geländeerhebungen etc.) unberücksichtigt bleiben.

■ **Topographic Openness (OPN):**

Dieser Index ist mit dem Sky View Factor vergleichbar, betrachtet jedoch nicht eine virtuelle Halbkugel oberhalb des mathematischen Horizontes, sondern ermittelt für einen Punkt bzw. eine Rasterzelle den durchschnittlichen Zenitwinkel (positive OPN) bzw. den durchschnittlichen Nadirwinkel (negative OPN) in einem bestimmten Radius in Grad (YOKOYAMA et al. 2002). Dadurch liegen die Ausgabewerte jeweils zwischen 0 und 180 Grad.

In der folgenden Tabelle 3 sind die verwendeten lokalen Parameter für die einzelnen Terrainindices angegeben. Ergänzend sind die verwendeten Bibliotheken und R-Pakete aufgeführt.

In der Abbildung 6 sind die verwendeten Terrainindices in jeweils einer Variante für den bereits erwähnten Ausschnitt aus dem Trainingsgebiet 1 dargestellt.

Tab. 3: Für die Modellierung verwendete Terrainindices.

Terrainindex	Kürzel	Parametrisierung*	Softwarebibliothek	R-Paket
Difference from mean Elevation	DME	Radius: 3 m / 5 m / 10 m	WhiteboxTools (LINDSAY 2016)	whitebox (WU & BROWN 2022)
Laplace of Gauss Filter	LGF	sigma**: 0,75 m / 1,00 m		
Impoundment Size Index	ISI	Dammlänge: 3 m / 5 m / 7 m		
Sky View Factor	SVF	Radius: 3 m / 5 m / 7 m	SAGA GIS (CONRAD et al. 2015)	RSAGA (BRENNING et al. 2022)
Topographic Openness	OPN	Radius: 1 m / 5 m / 10 m		

* Die Parameter werden in den jeweiligen Funktionen in Pixeln angegeben. Im vorliegenden Fall gilt: 1 px $\triangleq$ 1 m (DGM1).
 ** Beim Laplace of Gauss Filter wird nicht der Radius, sondern die Standardabweichung in Pixeln angegeben. Hieraus wird der Radius dann automatisch ermittelt.

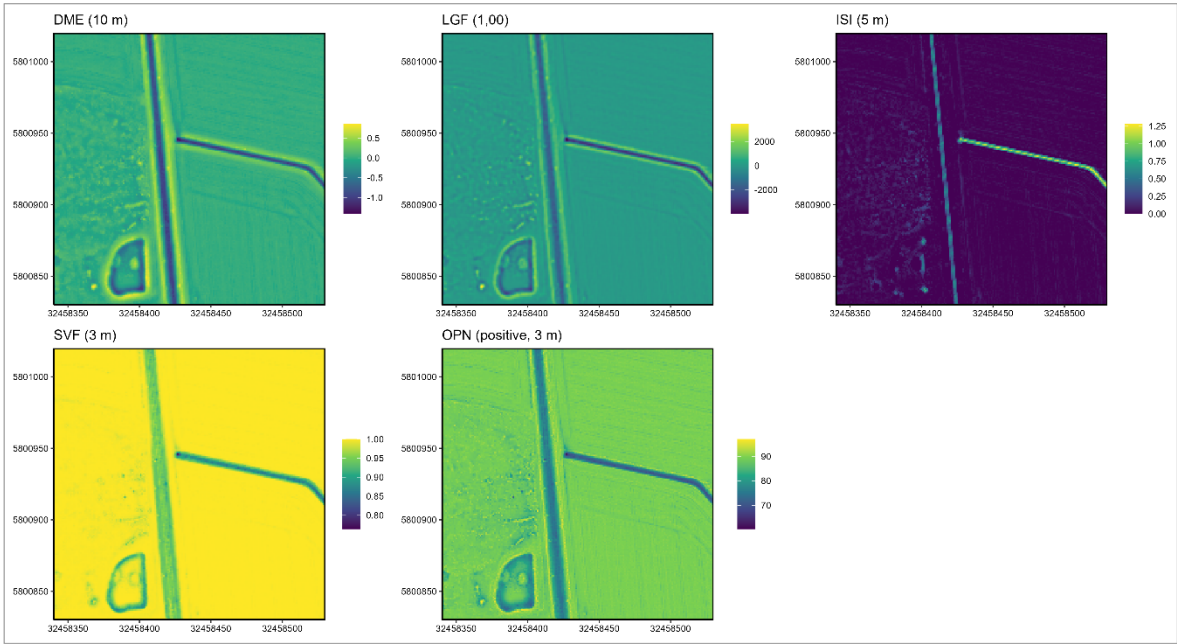


Abb. 6: Terrainindices für einen Ausschnitt aus dem Trainingsgebiet 1 (Wimmerbach). Abkürzungen und Parametrisierung gemäß Tabelle 2.

3.3. Pre-Processing der Trainingsdatensätze

Für die Zielvariable „Graben“ wurden im weiteren Verlauf lediglich zwei Ausprägungen vergeben: Eine Rasterzelle, welche zu einem Graben gehört, erhielt den Wert „1“; einer Rasterzelle außerhalb eines Grabens wurde der Wert „0“ zugewiesen. Im Zuge der Modellerstellung zeigte sich, dass die besten Ergebnisse erreicht werden konnten, wenn lediglich die Rasterzellen aus dem unmittelbaren Bereich der Grabensohlen als tatsächliche „Grabenpixel“ in das Modell eingingen. Um dies zu erreichen, wurden die im Vektorformat vorliegenden Linien der digitalisierten Grabenverläufe mit einem Buffer von 0,1 m versehen. Anschließend wurden die Rasterzellen, welche die zuvor erzeugten Polygone berührten, als „Graben“ definiert und mit einer „1“ versehen. Alle anderen Rasterzellen erhielten den Wert „0“ (s. Abb. 7, links).

Durch das Vorgehen mit einem vergleichsweise sehr kleinen Buffer wurden allerdings – je nach Breite der jeweiligen Gräben – auch Rasterzellen mit einer „0“ versehen, die tatsächlich zu einem Graben gehören. Um einen möglichst guten Trainingsdatensatz zu erhalten, wurden daher die Rasterzellen, welche sich in einem Abstand von unter 3,5 m zur Zentrallinie befanden,

nicht in den Trainingsdatensätzen berücksichtigt (s. Abb. 7, Mitte). Des Weiteren wurden die Häufigkeitsverteilungen sämtlicher Indices innerhalb der Trainingsdaten betrachtet, um etwaige nicht plausible Datenpunkte zu erkennen und ggf. zu entfernen. Im Ergebnis wurden die Rasterzellen vom Training ausgeschlossen, welche als „Graben“ gekennzeichnet waren und einen ISI-Wert (5 m) von 0 oder einen LGF-Wert (1,00 m) von < 0 aufwiesen. Dies betraf allerdings nur vergleichsweise sehr wenige Rasterzellen.

Nach den vorgenannten Bearbeitungsschritten umfassten die vorläufigen Trainingsdatensätze der einzelnen Trainingsgebiete noch zwischen 3.650.727 und 3.926.765 der ursprünglich jeweils 4.000.000 Rasterzellen. Hierbei gehörten lediglich zwischen rund 0,4 % und 2,0 % der Rasterzellen zu einem Graben. Bei diesen Verhältnissen wurden die Modelle zumeist zu stark auf die „Nicht-Graben-Pixel“ trainiert, so dass sie unempfindlicher bei der Erkennung von Grabenstrukturen wurden. Daher wurden die Trainingsdatensätze ausbalanciert, indem die Anzahl der „Nicht-Graben-Pixel“ randomisiert auf die Anzahl der „Graben-Pixel“ reduziert wurde (s. Abb. 7, rechts). Die endgültigen Trainingsdatensätze enthielten – je nach Trainingsgebiet – zwischen 29.618 und 142.826 Rasterzellen.

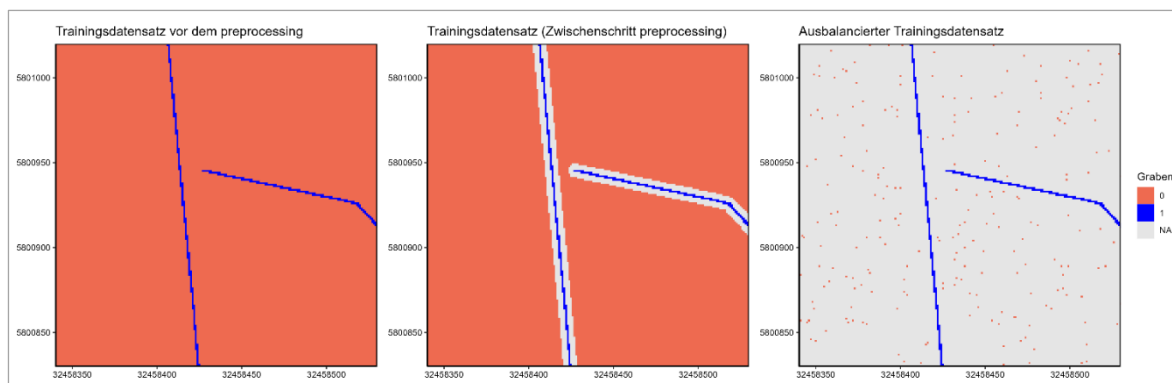


Abb. 7: Trainingsdatensatz für einen Ausschnitt aus dem Trainingsgebiet 1 (Wimmerbach) vor dem Pre-Processing (links), nach Ausschluss der Böschungsbereiche (Mitte) und nach Ausbalancierung (rechts). Die grau dargestellten Rasterzellen wurden nicht für das Modelltraining berücksichtigt.

3.4. Aufbau von Random-Forest-Modellen und Auswahl des besten Modells

Für jede Rasterzelle der Trainingsdatensätze lagen nach Durchführung der vorgenannten Schritte sowohl die Ergebniswerte der einzelnen Terrainindices als auch die Information, ob es sich bei der Rasterzelle um einen Graben handelt oder nicht, vor. Für das anschließende Modelltraining stellte die Grabeninformation die sogenannte Zielvariable dar, die Terrainindices wurden als Prädiktorvariablen eingesetzt.

Bei der Erstellung von Random-Forest-Modellen wird eine hohe Anzahl an Entscheidungsbäumen erzeugt. Für jeden dieser Bäume wird zunächst ein individueller Trainingsdatensatz generiert, indem aus dem „originalen“ Trainingsdatensatz eine zufällige Stichprobe mittels „Ziehen und Zurücklegen“ entnommen wird (BREIMAN 2001). Diese Methode wird auch als „bagging“ (ein Akronym für „bootstrap aggregating“) bezeichnet (BREIMAN 1996). Für die einzelnen Trainingsdatensätze wird nun – ebenfalls zufällig – eine begrenzte Anzahl an Prädiktorvariablen ausgewählt („random feature selection“), anhand derer letztlich ein Entscheidungsbaum erzeugt und vollständig bzw. bis zu einem festgelegten Grad ausgebaut wird (BREIMAN 2001, HASTIE et al. 2009). In der Anwendung des Modells durchlaufen unbekannte Datenpunkte sämtliche Entscheidungsbäume, wobei jeder Baum ein Einzelergebnis für den jeweiligen Datenpunkt erzeugt. Schließlich werden diese Einzelergebnisse zu einer Gesamtantwort für den unbekannten Datenpunkt aggregiert (HASTIE et al. 2009).

Im vorliegenden Fall wurden für die einzelnen Trainingsgebiete jeweils mehrere binäre Random-Forest-Klassifikationsmodelle (Graben = 1, kein Graben = 0) mit dem R-Paket ranger (WRIGHT & ZIEGLER 2017) erstellt. Als erste Zielgröße zur Einschätzung der Modellgüte und zur Auswahl der Prädiktorvariablen (hier: Terrainindices) diente der sogenannte „out-of-bag error“ (OOB error). Dieser wird während des Modelltrainings ermittelt, indem zunächst die durch die Methode des baggings unberücksichtigten Datenpunkte die Entscheidungsbäume durchlaufen, bei denen sie „out-of-bag“ waren. Die Ergebnisse werden anschließend mit den tatsächlichen Datenpunkten abgeglichen und zu einem Gesamtfehler aggregiert. Hierbei zeigte sich, dass – je nach verwendeter Hyperparameter-

kombination – die Modelle aus den Trainingsgebieten 1 (Wimmerbach) und 10 (Alte Kapelle) zumeist die niedrigsten Fehlerwerte und somit „intern“ die jeweils beste Vorhersagequalität aufwiesen.

Anschließend wurde geprüft, wie gut die einzelnen Modelle die Daten aus den jeweils anderen Trainingsgebieten vorhersagen können. Hierzu wurden die Modelle sowohl auf die balancierten als auch auf die unbalancierten Trainingsdatensätze angewendet und die Ergebnisse unter anderem mit dem R-Paket cvms (OLSEN & ZACHARIAE 2024) ausgewertet. Als Zielgröße für eine Einschätzung der Modellgüte wurde unter anderem der sogenannte Kappa-Index (auch: Cohen's Kappa) verwendet (COHEN 1960). Dieser Index korrigiert das Maß der Übereinstimmung zwischen den prognostizierten und den tatsächlichen Werten mit dem Anteil, welcher bei einer rein zufälligen Prognose zu erwarten wäre. Die Spanne der möglichen Werte des Kappa-Indices liegt zwischen -1 und 1, wobei eine vollständige Übereinstimmung von tatsächlichen Daten und prognostizierten Daten zu einem Wert von 1 führt. Ein Kappa-Index von 0 zeigt, dass die Prognose nicht besser ausfällt als eine rein zufällige Verteilung, ein Wert von -1 zeigt eine vollständig falsche Prognose an (BEN-DAVID 2008). Nähere Angaben zur Berechnung des Kappa-Indices finden sich im Abschnitt 6.2.

Bei der Anwendung auf die balancierten Trainingsdatensätze ergaben sich für die einzelnen Modelle – je nach Parametrisierung – zumeist mittlere Kappa-Indices von $> 0,8$ und damit sehr hohe Übereinstimmungen (LANDIS & KOCH 1977). Hierbei zeigte jedoch vor allem das Modell aus dem Trainingsgebiet 4 (Hüsedede) vereinzelt sehr geringe Werte, da in bestimmten Bereichen die vorhandenen Gräben nicht erkannt werden konnten. Auch das Modell aus dem Trainingsgebiet 10 (Alte Kapelle) wies vereinzelt geringe Werte auf. Bei der Anwendung auf die unbalancierten Trainingsdatensätze ergab sich zum Teil eine andere Verteilung der Kappa-Indices. Hier wurden – ebenfalls abhängig von der konkreten Parametrisierung – zumeist mittlere Werte zwischen rund 0,4 und rund 0,7 ermittelt, was gemäß LANDIS & KOCH (1977) mittleren bis hohen Übereinstimmungen entspricht. Insgesamt wiesen bei der Anwendung auf die unbalancierten Trainingsdatensätze die Modelle, welche die tatsächlichen Gräben gut erkennen konnten, tendenziell geringere Werte auf. Dies ist vor allem darauf zurückzuführen, dass diese Modelle häufiger auch Rasterzellen

als Gräben auswiesen, welche im Trainingsdatensatz keinem Graben zugeordnet worden waren (sogenannte „false positives“). Im vorliegenden Fall ist dies jedoch nicht per se problematisch, da es durchaus vorteilhaft sein kann, wenn das Modell beispielsweise einen Graben erkennt, welcher zuvor nicht eindeutig als solcher erkannt bzw. gekennzeichnet worden war.

Bei der Auswahl des am besten geeigneten Random-Forest-Modells wurden zum einen die rechnerisch ermittelten Güteparameter herangezogen. Parallel wurden die Ergebnisse zwischenzeitlich immer wieder im GIS visualisiert und begutachtet. Hierbei wurden auch Kombinationen von verschiedenen Modellen getestet. Insgesamt konnte festgestellt werden, dass die Modelle aus den Trainingsgebieten der flacheren Regionen häufig auch für stärker reliefierte Landschaften bessere Ergebnisse lieferten als die Modelle aus den Trainingsgebieten des Berg- und Hügellandes. Des Weiteren zeigten einige Modelle zum Teil „Ausfälle“, indem in bestimmten Bereichen die Grabenstrukturen nur unzureichend erkannt wurden. Andererseits wiesen einige Modelle nahezu durchgehend zu viele Rasterzellen als Graben aus. Schließlich wurde das Modell aus dem Trainingsgebiet 1 (Wimmerbach) für die weitere Bearbeitung ausgewählt. Zum einen wurden die tatsächlich vorhandenen Grabenstrukturen in den unter-

schiedlichen Gebieten von dem Modell insgesamt gut bis sehr gut erkannt. Zum anderen wurden im Vergleich zu den anderen Modellen, welche ebenfalls eine gute Wiedererkennung der tatsächlichen Gräben aufwiesen, vergleichsweise weniger Strukturen fälschlicherweise als Graben ausgewiesen.

3.5. Parametrisierung und Interpretation des finalen Random-Forest-Modells

Das Random-Forest-Modell aus dem Trainingsgebiet 1 (Wimmerbach) wurde für die Anwendung auf die Landesfläche Niedersachsens weiter optimiert und die Ergebnisse zwischenzeitlich immer wieder visuell kontrolliert. Auf die verwendeten Prädiktorvariablen wurde im Abschnitt 3.2 bereits eingegangen. Die endgültige Parametrisierung geht aus der Tabelle 4 hervor. Im Gegensatz zu den bisher erstellten binären Klassifikationsmodellen wurde schließlich ein Regressionsmodell erstellt, so dass anstelle einer Zugehörigkeitsklasse (Graben bzw. Nicht-Graben) eine „Grabenwahrscheinlichkeit“ für jede Rasterzelle ausgewiesen wurde. Bei dem Aufbau der Bäume wurde als „Unreinheitsmaß“ die Varianz des jeweiligen Datensatzes herangezogen, welche innerhalb des Modells an jedem Entscheidungsknoten minimiert wird.

Tab. 4: Parametrisierung des finalen Random-Forest-Modells.

Parameter	Ausprägung	Bemerkung
type	regression	Zielvariable wird numerisch behandelt (Graben = 1, Nicht-Graben = 0), die Ausgabe ist somit eine Wahrscheinlichkeit zwischen 0 und 1, dass die jeweilige Rasterzelle zu einem Graben gehört.
trees	500	Anzahl der erzeugten Entscheidungsbäume im Modell.
mtry	3	Anzahl der zufällig ausgewählten Prädiktorvariablen für die Erstellung eines Entscheidungsbaums.
min_n	5	Mindestanzahl an Datenpunkten am letzten Entscheidungsknoten.
max_depth	–	Keine maximale Tiefe der Bäume definiert; alle Entscheidungsbäume werden bis zum min_n ausgebaut.

Die Verringerung der „Unreinheit“ des Trainingsdatensatzes innerhalb des Modells kann zudem dazu verwendet werden, die Bedeutung der einzelnen Prädiktorvariablen zu quantifizieren: Die sogenannte Variablenwichtigkeit („feature importance“) beschreibt, wie stark die jeweilige Variable die Unreinheit der Trainingsda-

tensätze minimiert bzw. den Informationsgewinn steigert, wenn sie zum Teilen der Daten verwendet wird (HASTIE et al. 2009, WRIGHT & ZIEGLER 2017). Für die Extraktion der Variablenwichtigkeit aus den Modellen wurde das R-Paket vip (GREENWELL & BOEHMKE 2020) verwendet.

In der nachfolgenden Abbildung 8 ist die relative Variablenwichtigkeit für das finale Random-Forest-Modell dargestellt.

Demnach wurde im Rahmen des Trainings vor allem durch den Einsatz von jeweils einer Variante der Indices „Laplace of Gauss Filter“ (LGF) sowie „Difference from mean Elevation“ (DME)

die Unreinheit der Datensätze verringert. Auch jeweils eine Variante des „Sky View Factors“ (SVF) sowie der „Openness“ (OPN) hatte einen vergleichsweise großen Effekt beim Teilen der Daten. Bei den weiteren Prädiktorvariablen fällt der Einfluss anschließend mehr oder weniger kontinuierlich ab.

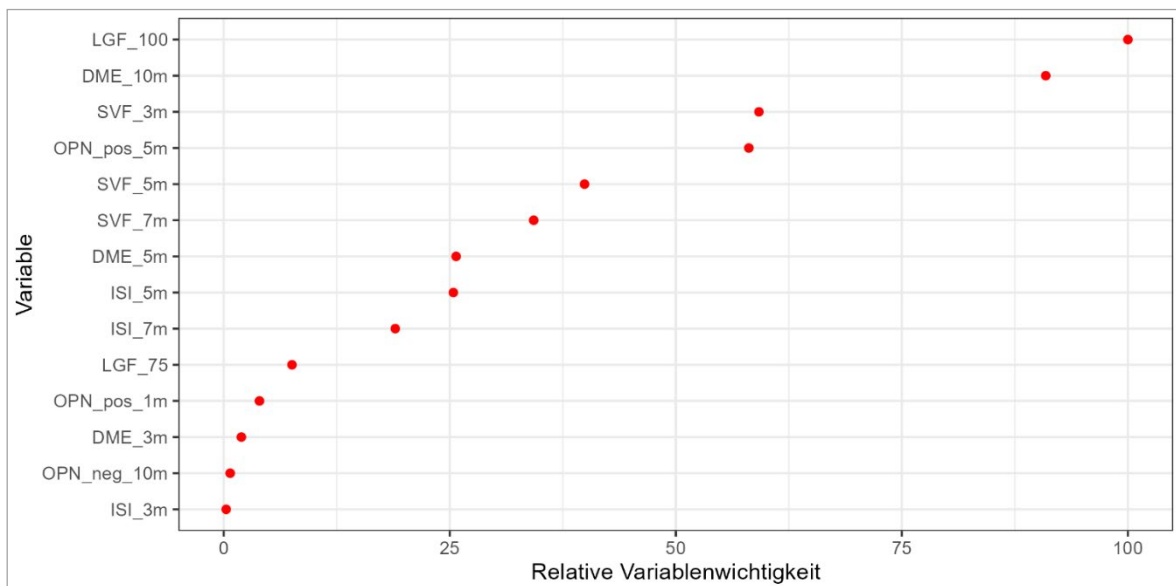


Abb. 8: Relative Variablenwichtigkeit des finalen Random-Forest-Modells.

3.6. Zwischenergebnis nach Anwendung des finalen Random-Forest-Modells

Die Abbildung 9 zeigt – beispielhaft für einen Ausschnitt des Trainingsgebietes 10 (Alte Kapelle) – die Prognoseergebnisse des finalen Random-Forest-Modells. Für jede Rasterzelle ist eine „Grabenwahrscheinlichkeit“ zwischen 0 und 1 angegeben. Ergänzend sind das zugehörige Digitale Orthophoto, die Schummerung des DGM1 sowie die zuvor manuell digitalisierten Gräben dargestellt.

Demnach weisen die Rasterzellen der zuvor digitalisierten Gräben allgemein hohe Grabenwahrscheinlichkeiten auf. Die sonstigen Rasterzellen sind zumeist durch geringe bis sehr geringe Werte charakterisiert. Teilweise können etwas erhöhte Werte im Bereich sogenannter Gruppen innerhalb der bewirtschafteten Flächen sowie in den Randbereichen breiterer Fließgewässer auftreten. Insgesamt ist die

Prognose als sehr gut einzuschätzen. Hierbei ist allerdings zu berücksichtigen, dass das Trainingsgebiet 10 (Alte Kapelle) in der Wesermarsch liegt und eine sehr niedrige topographische Variabilität aufweist. Wie bereits erwähnt, ist davon auszugehen, dass die Ungenauigkeit der Prognose tendenziell mit der Komplexität der Topographie sowie mit einem zunehmenden Baumbestand steigt.

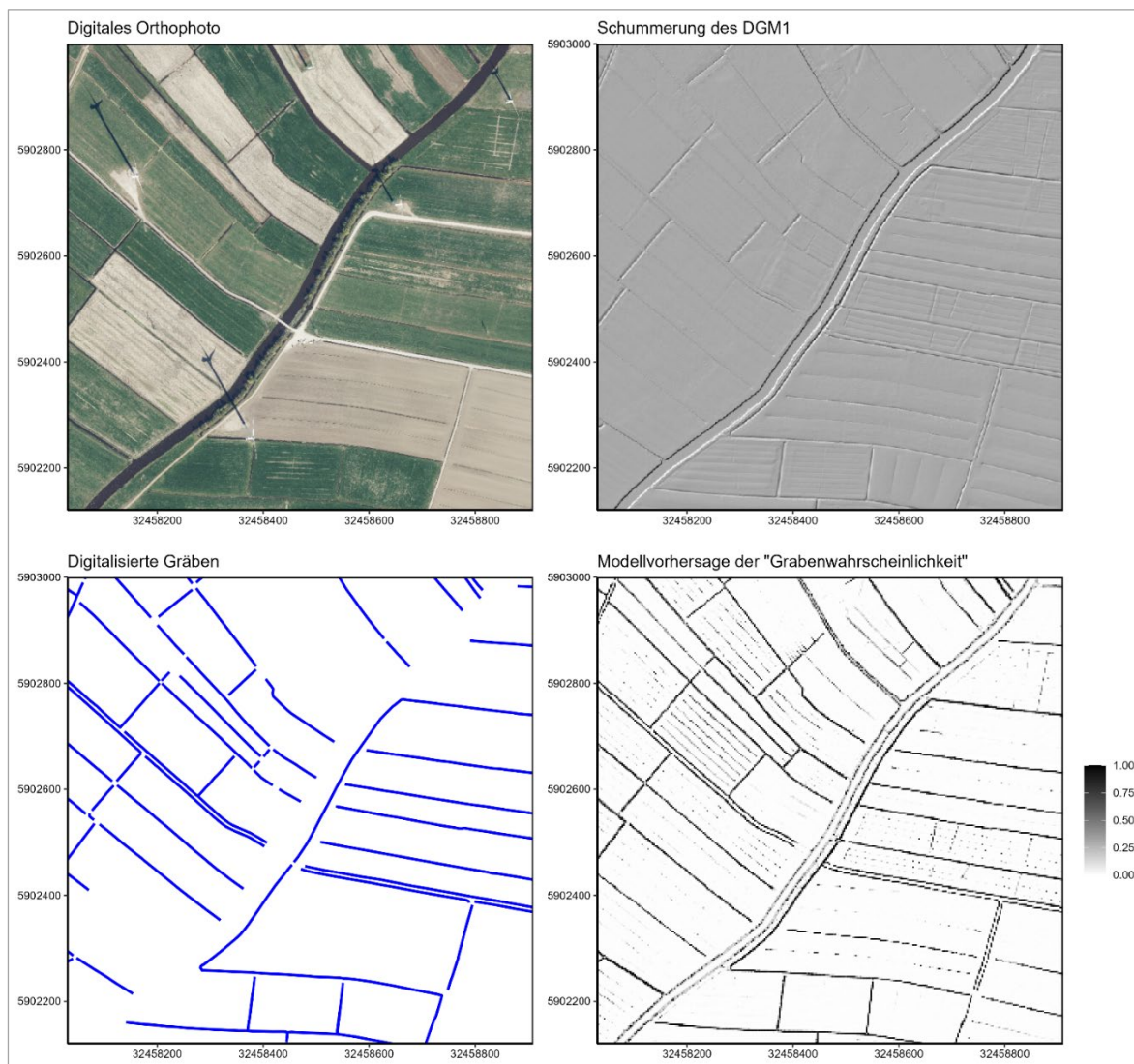


Abb. 9: Ausschnitt aus dem Trainingsgebiet 10 (Alte Kapelle) mit dem zugehörigen Digitalen Orthophoto, der DGM1-Schumierung, den manuell digitalisierten Grabenverläufen sowie dem Ergebnistraster nach Anwendung des finalen Random-Forest-Modells.

Für die Abbildung 10 wurde daher ein Ausschnitt des Trainingsgebietes 4 (Hüsedede) ausgewählt, welcher im Übergangsbereich vom Bergvorland zum Bergland (s. auch Abb. 3) liegt. Während im östlichen Bereich des Ausschnittes landwirtschaftliche Nutzflächen mit geringer Neigung dominieren, ist der südwestliche Bereich bereits durch die bewaldeten Höhenzüge des Wiehengebirges gekennzeichnet. Im Norden des Kartenausschnittes befinden sich zudem Siedlungsflächen.

Aus der Abbildung 10 geht hervor, dass – analog zum Trainingsgebiet 10 – die Rasterzellen der Entwässerungsgräben insgesamt hohe Werte aufweisen. Allerdings sind vor allem im

Bereich der bewaldeten Flächen mit stärkerer Reliefenergie ebenfalls viele Rasterzellen mit erhöhten Wahrscheinlichkeitswerten vorhanden, die zu keinem Graben gehören. Dies betrifft etwa Geländesprünge oder -einschnitte, kleinere natürliche Gewässer sowie (ehemalige) Hohlwege, die sich in das Gelände eingeschnitten haben.

Um die vorgenannten „Fehler“ zu korrigieren, wurde daher ein mehrstufiges Post-Processing entwickelt, um möglichst viele dieser Rasterzellen, welche zu falsch positiven Ergebnissen führen würden, auszuschließen. Hierauf wird in den folgenden Abschnitten 4 und 5 näher eingegangen.

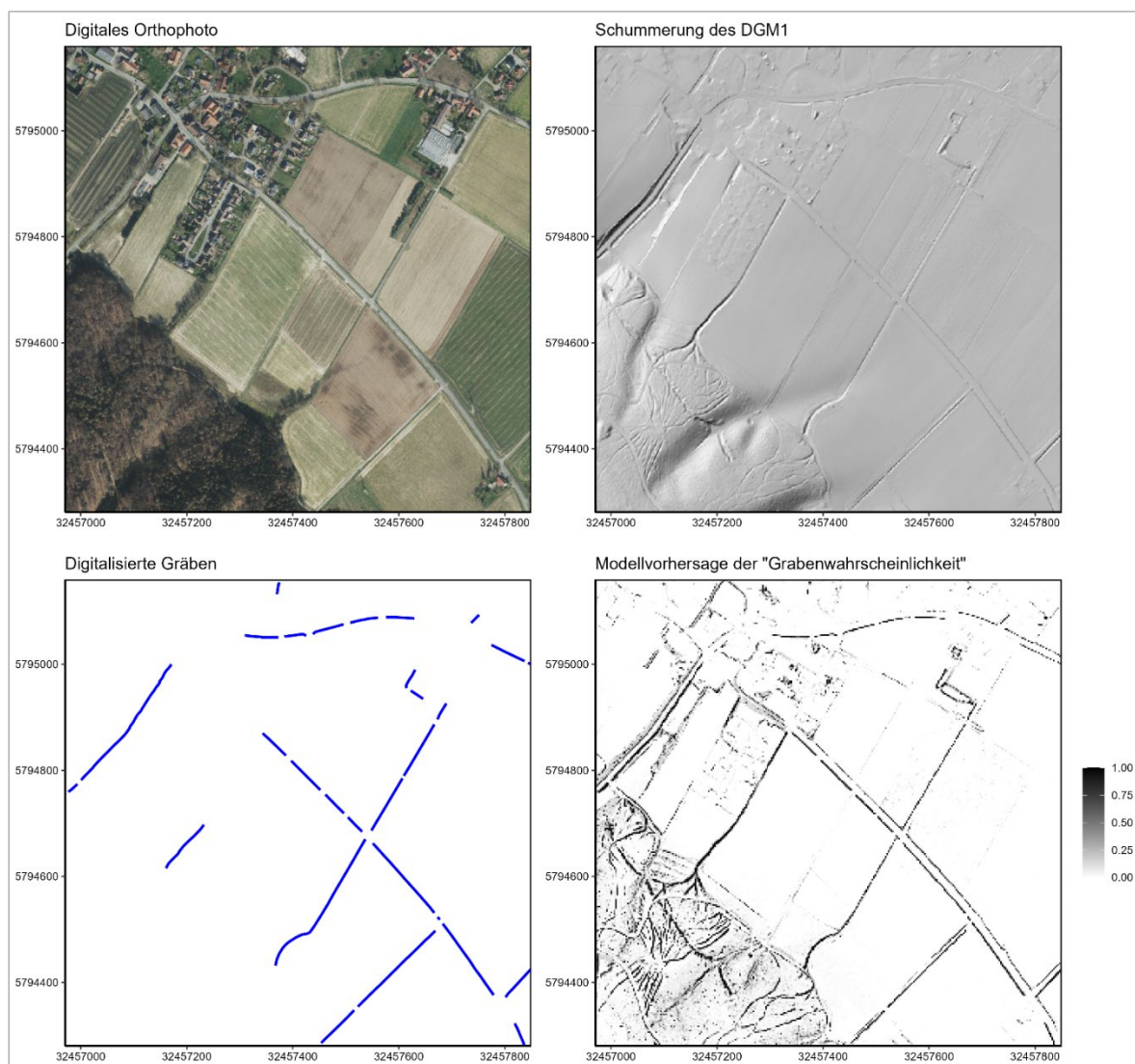


Abb. 10: Ausschnitt aus dem Trainingsgebiet 4 (Hüsedede) mit dem zugehörigen Digitalen Orthophoto, der DGM1-Schummierung, den manuell digitalisierten Grabenverläufen sowie dem Ergebnistraster nach Anwendung des finalen Random-Forest-Modells.

4. Clustern von Grabenstrukturen

In diesem Abschnitt wird die Optimierung des Zwischenergebnisses aus dem Random-Forest-Modell mit verschiedenen Verfahren beschrieben. Die Karte der Grabenwahrscheinlichkeit wird hierbei in eine Karte mit „Grabenclustern“ überführt.

4.1. Ausschlusskriterien

Zunächst wurden auf Basis der im vorhergehenden Abschnitt 3.6 angesprochenen Befunde verschiedene Ausschlusskriterien definiert. Hierbei wurde davon ausgegangen, dass Entwässerungsgräben unter bestimmten Bedingungen nicht oder nur in vernachlässigbarem Umfang vorhanden sind. Sämtlichen Rasterzellen, welche in den nachfolgend aufgeführten Bereichen liegen, wurde daher der Wert 0 zugewiesen.

■ Feldblöcke:

Landwirtschaftliche Entwässerungsgräben befinden sich – mit Ausnahme von flachen Gruppen zur Aufnahme von oberflächlich abfließendem Wasser – in der Regel nicht innerhalb der tatsächlichen Nutzflächen, sondern verlaufen seitlich in unmittelbarer Nähe zu den Schlägen bzw. Feldblöcken. Um sicherzustellen, dass direkt an den Grenzen verlaufende Gräben bei der Modellierung erhalten blieben, wurden die Feldblockgeometrien (SLA 2022) zuvor mit einem inneren Buffer von 3 m versehen.

■ Flächenhafte Gewässer:

An den Randbereichen größerer Fließ- und Stillgewässer wurden zum Teil erhöhte Grabenwahrscheinlichkeiten ausgewiesen. Um diese Bereiche sicher auszuschließen, wurden die Gewässer, die über eine Flächeninformation verfügen (LGLN 2024c), zuvor mit einem Buffer von 5 m versehen.

■ Konvexe Geländeformen und stark geneigte Hänge:

Die tendenziell wenigen im Bereich des Berg- und Hügellandes vorhandenen Entwässerungsgräben sind häufig entweder Straßenseitengräben oder befinden sich seitlich von vergleichsweise gering geneigten landwirtschaftlichen Flächen. Daher wurden konvexe Geländeformen wie etwa Kuppenbereiche sowie stark geneigte Hänge von der Modellierung ausgeschlossen. Die Berechnung der Geländeformen („Geomorphons“) erfolgte auf Basis des DGM1 mit dem Paket whitebox (Wu & BROWN 2022).

■ Tiefe Geländeeinschnitte:

Für den Ausschluss tieferer Geländeeinschnitte – vor allem im Bereich natürlicher Gewässer im Bergland – wurde zunächst der Impoundment Size Index (s. Abschnitt 3.2) mit einer Dammlänge von 10 m ermittelt. Für die Abgrenzung wurde nach visueller Prüfung und dem Abgleich mit vorhandenen Entwässerungsgräben ein Schwellenwert von 2,8 m für die Dammhöhe gewählt.

■ Küstendüne:

Die Dünen der Nordseeinseln wiesen aufgrund ihrer Topographie viele Rasterzellen mit erhöhten Grabenwahrscheinlichkeiten auf. Da es sich hierbei nicht um Entwässerungsgräben handelt, konnten die Bereiche von der Modellierung ausgeschlossen werden. Die Geometrien wurden der BK50 (GEHRT et al. 2021) entnommen.

Je nach naturräumlichen Gegebenheiten hatte die Anwendung der vorgenannten Ausschlusskriterien unterschiedlich starke Auswirkungen auf das Ergebnis. In der Abbildung 11 sind die Zwischenschritte erneut beispielhaft für die beiden Ausschnitte aus den Trainingsgebieten 10 und 4 visualisiert. In der Wesermarsch kommt vor allem der Ausschluss der flächenhaften Gewässer sowie der Feldblöcke zur Geltung, während sich im Berg- und Hügelland vor allem der Ausschluss bestimmter Geländeformen sowie tiefer Geländeeinschnitte bemerkbar macht.

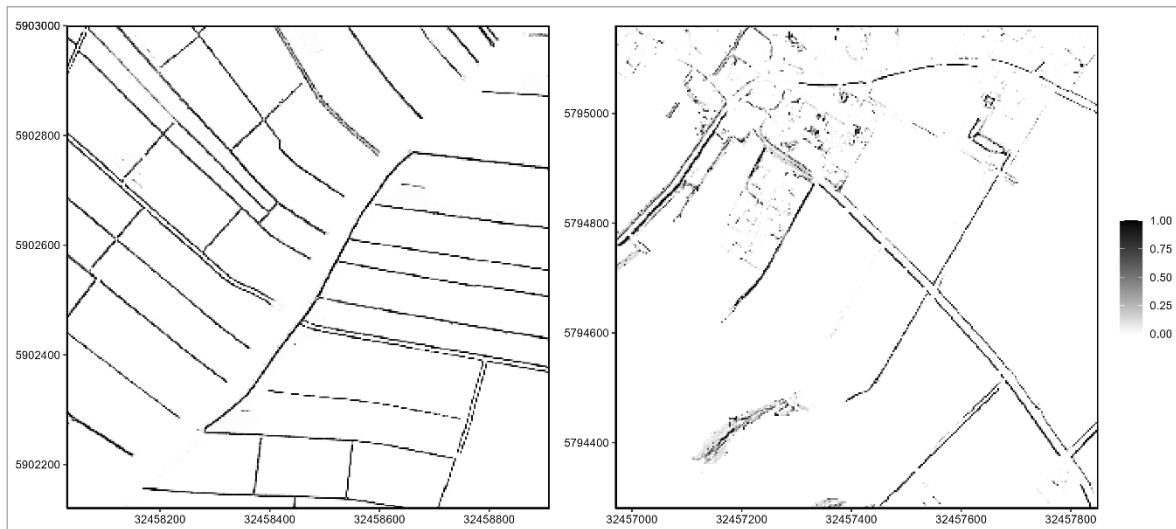


Abb. 11: Zwischenergebnis nach Anwendung der Ausschlusskriterien. Links: Ausschnitt aus dem Trainingsgebiet 10 (Alte Kapelle), rechts: Ausschnitt aus dem Trainingsgebiet 4 (Hüsedede).

4.2. Entrauschen und Clustern

Um weiteres „Rauschen“ zu reduzieren, wurden die folgenden Schritte nacheinander durchgeführt.

■ Bilateraler Filter:

Bilaterales Filtern kann als Weichzeichner eingesetzt werden, wobei im Vergleich zu anderen Filtermethoden Kanten besser erhalten bleiben (TOMASI & MANDUCHI 1998). Im vorliegenden Fall werden somit zum Beispiel einzelne Rasterzellen, welche eine hohe Grabenwahrscheinlichkeit aufweisen, jedoch isoliert von anderen Rasterzellen mit ebenfalls hohen Werten liegen, herausgefiltert. Auf der anderen Seite werden Rasterzellen mit etwas geringen Wahrscheinlichkeiten „aufgewertet“, sofern sie zwischen Rasterzellen mit hohen Werten liegen.

■ Binarisation:

Um konkrete Grabenstrukturen zu erhalten, wurde die Zielvariable binarisiert, also in eine Klassenstruktur mit lediglich zwei Ausprägungen überführt. Hierfür wurde nach intensiver visueller Prüfung ein Schwellenwert von 0,4 gewählt. Rasterzellen mit einer Grabenwahrscheinlichkeit von $> 0,4$ erhielten den Wert 1, Rasterzellen mit einer Grabenwahrscheinlichkeit $\leq 0,4$ den Wert 0. Hierbei entstanden zusammenhängende Cluster an „Grabenpixeln“.

■ Trimmen:

Die im Zuge der Binarisation entstandenen Cluster wiesen zum Teil noch kleinere Unregelmäßigkeiten, Ausläufer oder seitliche Vorsprünge auf. Diese wurden mit Hilfe des Paketes *whitebox* (WU & BROWN 2022) entfernt. In diesem Zuge wurden auch sehr kleine Cluster herausgefiltert.

Die Abbildung 12 zeigt das Ergebnis nach Anwendung der vorgenannten drei Schritte.

Demnach entsprachen die resultierenden Zwischenergebnisse in den flachen Bereichen größtenteils den tatsächlichen Grabenverläufen. Zwar wurde vor allem im Bergland bzw. in Waldflächen sowie in Siedlungsgebieten das Rauschen reduziert, dennoch konnten auch hiermit vergleichsweise viele Strukturen, bei denen es sich nicht um Gräben handelt, nicht herausgefiltert werden. So waren nach wie vor einige natürliche Geländeeinschnitte, Hohlwege

sowie künstliche Geländesprünge im Seitenbereich von Straßen in den Ergebnisrastern vorhanden. „Einfache“ Anpassungen in den bisherigen Prozessschritten, wie etwa ein Anheben des Schwellenwertes bei der Binarisation (s. Abschnitt 4.2), führten zwar häufig zu einem Ausschluss dieser Strukturen. Allerdings wurden gleichzeitig stets auch einige der tatsächlich vorhandenen Gräben herausgefiltert. Daher wurde ein weiterer, auf Maschinellem Lernen basierender Filter entwickelt, der im folgenden Abschnitt 5 beschrieben wird.

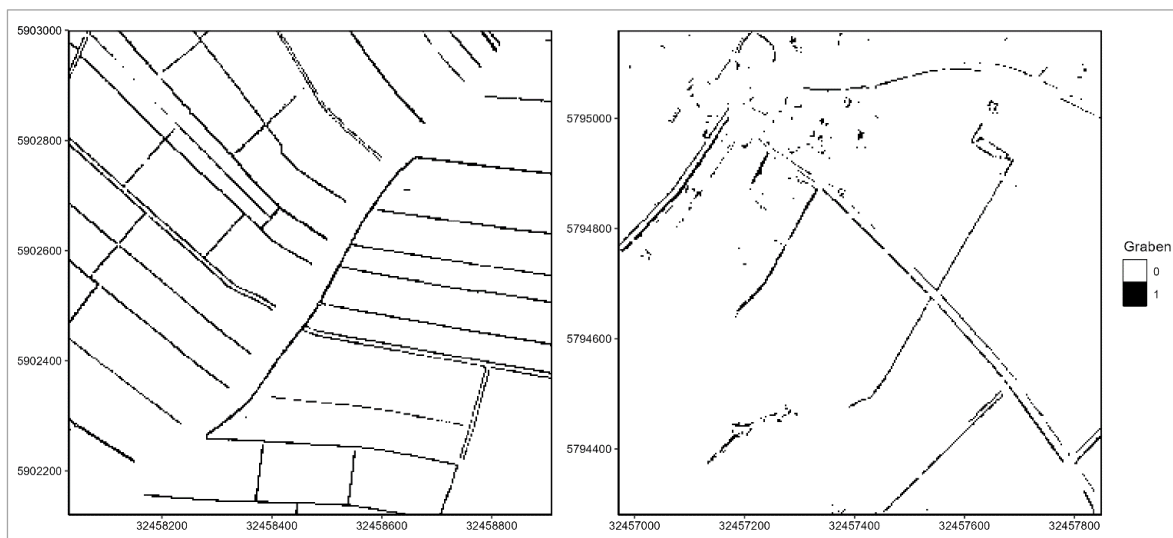


Abb. 12: Zwischenergebnis nach dem Entrauschen und Clustern. Links: Ausschnitt aus dem Trainingsgebiet 10 (Alte Kapelle), rechts: Ausschnitt aus dem Trainingsgebiet 4 (Hüsedede).

5. Filtern der Graben-Cluster

Bei der Erstellung des Random-Forest-Modells (s. Abschnitt 3) war im Wesentlichen davon ausgegangen worden, dass sich die einzelnen Rasterzellen eines Grabens durch bestimmte Eigenschaften (hier: Terrainindices) von anderen Rasterzellen unterscheiden. Für den im Folgenden beschriebenen Ansatz wurde angenommen, dass sich zusätzlich auch die aus vielen einzelnen Rasterzellen bestehenden „Graben-cluster“ von anderen Clustern unterscheiden lassen. Um möglichst viele Cluster, welche bislang fälschlicherweise als Entwässerungsgraben erkannt wurden, herauszufiltern und gleichzeitig möglichst viele der korrekt erkannten Grabencluster zu erhalten, wurde mit dem XGBoost-Algorithmus (CHEN & GUESTRIN 2016)

erneut eine Methode des Maschinellen Lernens genutzt. Das Ergebnis nach dem „Herausfiltern“ ist eine Karte mit Clustern, welche jeweils eines der beiden Label „Grabenstruktur wahrscheinlich“ oder „Grabenstruktur möglich“ aufweisen.

5.1. Anreichern der Cluster mit weiteren Informationen

Während der visuellen Kontrollen der Zwischenergebnisse wurden neben der bereits im Abschnitt 2.2 bzw. Abbildung 1d erwähnten linearen Form weitere Parameter identifiziert, welche für eine Abgrenzung von Entwässerungsgräben gegenüber anderen Strukturen in Frage kommen könnten. Die schließlich ausgewählten Parameter lassen sich grob in die drei Gruppen „Form“, „Naturraum“ und „Lage“ einteilen und werden nachfolgend näher erläutert. Für die Berechnungen kamen unter anderem die R-Pakete terra (HIJMANS 2023), tidyverse (WICKHAM et al. 2019) und whitebox (WU & BROWN 2022) zum Einsatz.

Aus der Gruppe „Form“ wurden die folgenden Parameter berechnet:

■ Fläche (A) [m²]:

Für jedes Cluster wurde die Anzahl an zusammenhängenden Rasterzellen ermittelt. Hierbei wurden jeweils alle acht möglichen Richtungen (N, NO, O, SO, S, SW, W, NW) berücksichtigt. Die Anzahl der Rasterzellen entspricht im vorliegenden Fall der Fläche des Clusters in m².

■ Maximale Distanz ($D_{\max}$) [m]:

Innerhalb eines jeden Clusters wurde für jedes mögliche Paar an Rasterzellen die Distanz zwischen beiden Zellen ermittelt. Der Maximalwert dieser Entfernungen gibt einen Hinweis auf die räumliche Ausdehnung des Clusters.

■ Kompaktheit (K) [–]:

Ergänzend wurde ein Kompaktheitsindex für jedes Cluster berechnet. Hierbei wurde die Fläche eines Clusters (s. o.) in das Verhältnis zu der Fläche eines Kreises gesetzt, dessen Durchmesser der maximalen Distanz (s. o.) entspricht:

$$K = \frac{A}{\pi \left(\frac{D_{\max}}{2}\right)^2}$$

■ Mittlere Breite (B_{mean}) [m]:

Zunächst wurde für jede Rasterzelle i die Anzahl n der angrenzenden Zellen separat für jede der 8 möglichen Richtungen (s. o.) bestimmt. Anschließend wurden die Werte für die jeweils gegenüberliegenden Richtungen addiert, so dass sich vier Werte für

die zusammenhängenden Rasterzellen in den Richtungen Nord-Süd, Ost-West, Nordost-Südwest sowie Nordwest-Südost ergaben. Unter zusätzlicher Berücksichtigung der jeweils betrachteten Rasterzelle ergab sich in Anlehnung an CAZORZI et al. (2013) die Breite B des Clusters an der Rasterzelle i aus dem Minimalwert dieser vier Einzelwerte zu:

$$B_i = \min \left((n_{i,N} + n_{i,S}), (n_{i,O} + n_{i,W}), (n_{i,NO} + n_{i,SW}), (n_{i,NW} + n_{i,SO}) \right) + 1$$

Schließlich wurden die vorgenannten Minimumwerte über alle k Rasterzellen des Clusters gemittelt, so dass sich die mittlere Breite B_{mean} ergab:

$$B_{\text{mean}} = \frac{\sum_{i=1}^k B_i}{k}$$

■ Länge (L) [m]:

Als Index für die Länge eines Clusters wurde analog zu CAZORZI et al. (2013) das Verhältnis aus Fläche und mittlerer Breite herangezogen.

■ Tiefe (T_{mean} , $T_{0.5}$, $T_{0.75}$, T_{ratio}) [m]:

Die Tiefe eines Clusters an jeder Rasterzelle wurde grob aus dem bereits erwähnten Impoundment Size Index abgeleitet. Hierzu wurden Dammlängen von 5 m bzw. 7 m angesetzt (ISI 5, ISI 7). Um das jeweilige Cluster in seiner Gesamtheit zu charakterisieren, wurden der Mittelwert (T_{mean}) und der Median ($T_{0.5}$) über sämtliche Rasterzellen gebildet. Ergänzend wurde nach eingehender visueller Prüfung das 75%-Perzentil ($T_{0.75}$) ermittelt, welches in vielen Fällen in etwa der mittleren Tiefe der Tiefenlinie eines Grabens entsprach. Des Weiteren wurde das Verhältnis von $T_{0.75}$ aus ISI 5 und $T_{0.75}$ aus ISI 7 als weitere Variable (T_{ratio}) eingesetzt.

■ Volumen (V) [m³]:

Das Volumen eines Clusters wurde sehr vereinfacht ebenfalls anhand des Impoundment Size Indices abgeschätzt. Aufgrund der Rasterzellengröße von 1 m² wurden die ISI-5-Werte sämtlicher Rasterzellen hierzu aufsummiert.

Neben den Formparametern wurden die Cluster mit Informationen aus der Parametergruppe „Naturraum“ angereichert. Hierzu wurden die Daten aus LGLN (2024c) bzw. GEHRT et al. (2021) verwendet.

■ **Wald:**

Generell war im Zuge der Modellierung festgestellt worden, dass Flächen mit einem hohen Baumbestand ein erhöhtes Rauschen aufwiesen. Um dies in einem Modell zur Ergebnisfilterung zu berücksichtigen, wurde den Clustern eine entsprechende binäre Information angefügt (Wald = 1, kein Wald = 0). Um in unmittelbarer Nähe zu einer Waldfläche liegende Cluster nicht fälschlicherweise mit dem Label „Wald“ zu versehen, wurden sämtliche Waldflächen zuvor mit einem inneren Buffer von 10 m versehen.

■ **Bodenregion:**

Die Bodenkarte 1 : 50.000 (BK50) nach GEHRT et al. (2021) weist insgesamt sechs verschiedene Bodenregionen aus: Küstenholozän (BR1), Flusslandschaften (BR2), Geest (BR3), Bergvorland (BR4), Bergland (BR5) und Mittelgebirge/Harz (BR6). Diese Regionen wurden den Clustern als weiterer Naturraumparameter angefügt. Die teilweise vorhandenen Polygone ohne bodenlandschaftliche Zuordnung waren zuvor jeweils der sie umgebenden Bodenregion bzw. der Bodenregion mit der längsten gemeinsamen Grenze zugeordnet worden.

Um einen möglichen Einfluss klimatischer bzw. topographischer Gegebenheiten mittelbar zu berücksichtigen, wurde zusätzlich ein Parameter aus der Gruppe „Lage“ ermittelt.

■ **Mittlere Höhenlage:**

Aus dem DGM1 wurden die Geländehöhen [NHN] über sämtliche Rasterzellen eines Clusters gemittelt. Die mittleren Rechts- und Hochwerte waren testweise ebenfalls berechnet worden, führten jedoch zu keiner Verbesserung der Ergebnisse.

5.2. Labeling von Clustern und Entwicklung von XGBoost-Modellen

Zunächst wurden sämtliche ermittelten Cluster für die gesamte Landesfläche mit den vorgenannten Daten angereichert. Anschließend wurden insgesamt acht Kacheln ausgewählt, in denen zuvor intensive Plausibilitätsprüfungen der Ergebnisse – am Bildschirm sowie im Gelände – durchgeführt worden waren. Für jede dieser Kacheln wurden die einzelnen Cluster manuell mit einer Grabeninformation gelabelt. Da eine eindeutige Einstufung selbst bei einer Überprüfung im Gelände nicht immer möglich war, wurde neben den Kategorien „J“ (Graben) und „N“ (kein Graben) in Einzelfällen auch die Kategorie „U“ (unklar) verwendet. Dies betraf beispielsweise vergleichsweise flache Gräben, welche zwar oberflächlich abfließendes Wasser aufnehmen und ableiten können, jedoch keine nennenswerte Drainagewirkung für den Boden bzw. unterirdisches Wasser aufweisen. Des Weiteren war in Einzelfällen die Unterscheidung zwischen Gräben und natürlichen Gewässern mit grabenähnlicher Struktur nicht eindeutig möglich.

Die Abbildung 13 zeigt Beispiele für Cluster bzw. Strukturen, die anhand einer Plausibilitätsprüfung im Gelände überprüft, mit einem entsprechenden Label versehen und später für das Modelltraining verwendet wurden.

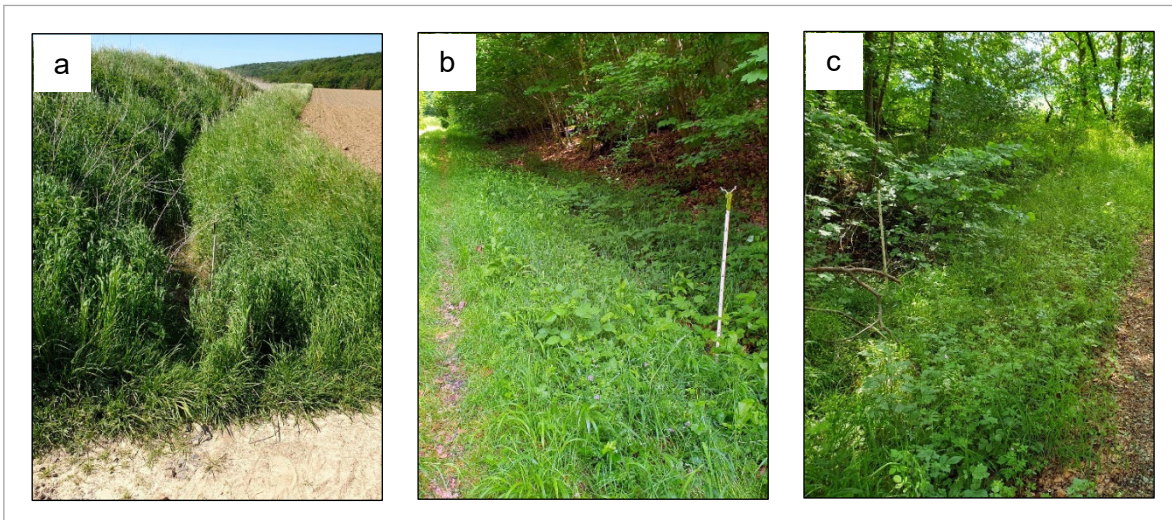


Abb. 13: Beispiele für Strukturen, die im Rahmen einer Plausibilitätsprüfung überprüft, mit einem Label versehen und später für das Modelltraining verwendet wurden. a: Graben, b: kein Graben, c: Grenzfall (flacher Graben). Fotos: M. Hoetmer.

Um einen umfangreicheren Trainingsdatensatz zu erhalten, wurden weitere 100 Ergebniskacheln randomisiert ausgewählt und die zugehörigen Cluster manuell mit einer Grabeninformation („J“, „N“, „U“; s. o.) versehen. Hierzu wurden die Cluster jeweils in einem GIS visuell hinsichtlich ihrer Lage und Form überprüft sowie mit Luftbildern und einer DGM1-Schummerung abgeglichen. Die insgesamt rund 10.000 manuell gelabelten Cluster dienten nachfolgend als Trainingsdatensatz für ein Gradient-Boosting-Tree-Ensemble.

Während bei Random Forests viele gleichwertige Bäume „nebeneinander“ erstellt werden (s. Abschnitt 3.4), besteht das Ensemble der Gradient Boosted Trees (GBT) aus „nacheinander“ erstellten Bäumen. Bei diesem iterativen Ansatz versucht jeder Baum, die „Fehler“ der vorherigen Bäume zu korrigieren und somit das Gesamtmodell zu optimieren („boosting“; FRIEDMAN 2001). Die Bäume sagen zudem nicht die Zielvariable selbst voraus, sondern den bisherigen Prognosefehler bzw. die Stärke der Abweichung der bisherigen Prognose von der Realität. Diese Abweichung wird somit im Laufe der Modellerstellung für jeden Datenpunkt möglichst minimiert (FRIEDMAN 2001). Des Weiteren werden die einzelnen Bäume tendenziell nicht so tief ausgebaut wie beispielsweise bei Random Forest und stellen somit jeder für sich vergleichsweise „schwache“ Modelle (sogenannte

„base learner“ bzw. „weak learner“) dar (FRIEDMAN 1999, 2001; HASTIE et al. 2009). Durch die iterative Vorgehensweise bilden viele dieser weak learner am Ende ein starkes Modellenensemble.

5.3. Parametrisierung und Modellinterpretation

Die Umsetzung erfolgte mit der Softwarebibliothek XGBoost (CHEN & GUESTRIN 2016) bzw. mit deren Implementierung in den R-Paketen xgboost (CHEN et al. 2024) und tidymodels (KUHN & WICKHAM 2020). Für die kategorialen Variablen (Wald, Bodenregion) wurden jeweils sogenannte dummy-Variablen erstellt, da die einzelnen Klassen keine natürliche Reihenfolge aufweisen. Die Hyperparametertuning erfolgte iterativ und mit dem Ziel, zunächst den bereits im Abschnitt 3.4 erwähnten Kappa-Index zu optimieren. Hierbei wurde mit einer stratifizierten 5-fold-cross-validation gearbeitet. Parallel dazu wurden – analog zu dem im Abschnitt 3.4 beschriebenen Vorgehen – auch immer wieder visuelle Kontrollen der Zwischenergebnisse durchgeführt, so dass nicht nur die berechneten Güteparameter für die Auswahl der „besten“ Hyperparameterzusammensetzung herangezogen wurden. Schließlich wurden zwei unterschiedliche Modelle trainiert, wobei für das erste Modell die Bodenregionen „Küstenholozän“,

„Flusslandschaften“, „Geest“ und „Bergvorland“ (BR1 bis BR4) zusammengefasst wurden und das zweite Modell für die Bodenregionen „Bergland“ und „Mittelgebirge/Harz“ (BR5 und BR6) entwickelt wurde. Das Verhältnis von Trainings- und Testdaten wurde mit 75:25 bzw. mit 80:20 gewählt. Die endgültige Parametrisierung der beiden Modelle geht aus der Tabelle 5 hervor. Bei der Anwendung auf die Testdaten wurden Kappa-Indices von 0,75 bzw. 0,76 erreicht, so dass gemäß LANDIS & KOCH (1977) jeweils eine hohe Übereinstimmung vorlag.

Die Abbildung 14 sowie die Abbildung 15 zeigen die relative Variablenwichtigkeit für die beiden XGBoost-Modelle. Bei dem Modell für die

Bodenregionen 1 bis 4 lieferten vor allem die Variablen „Tiefe ISI5 0.75“ und „Volumen ISI5“ die höchsten Informationsgewinne beim Teilen der Trainingsdaten. An dritter Stelle folgt die Variable „Kompaktheit“. Bei dem Modell für die Bodenregionen 5 und 6 ist die Kompaktheit mit deutlichem Abstand die Variable mit dem höchsten Informationsgewinn bei der Datenteilung. Anschließend folgen die Variablen „Tiefe ISI5 0.75“ sowie „Tiefe ratio 0.75“. Des Weiteren ist zu erkennen, dass bei beiden Modellen die mittlere Höhenlage sowie die maximale Distanz innerhalb der Cluster vergleichsweise wichtige Variablen darstellen.

Tab. 5: Parametrisierung der XGBoost-Modelle.

Parameter	Ausprägung		Bemerkung
	Bodenregionen 1 bis 4	Bodenregionen 5 und 6	
type	classification	classification	Zielvariable wird als Klasse behandelt, die Ausgabe entspricht der Klasse, welche einem unbekannten Datenpunkt überwiegend zugewiesen wird.
trees	125	125	Anzahl der erzeugten Bäume im Modell.
mtry	7	7	Anzahl der zufällig ausgewählten Prädiktorvariablen für die Erstellung eines Baums.
min_n	5	5	Mindestanzahl an Datenpunkten am letzten Entscheidungsknoten.
tree depth	4	4	Alle Bäume werden maximal bis zu dieser festgelegten Tiefe bzw. bis zum min_n ausgebaut.
learn_rate	0.10	0.05	Parameter, welcher den Einfluss des jeweils aktuellen Baumes auf das bisherige Modellergebnis begrenzt.
sample_size	0.75	0.70	Anteil der zufällig für jeden Baum gezogenen Datenpunkte aus dem Trainingsdatensatz.

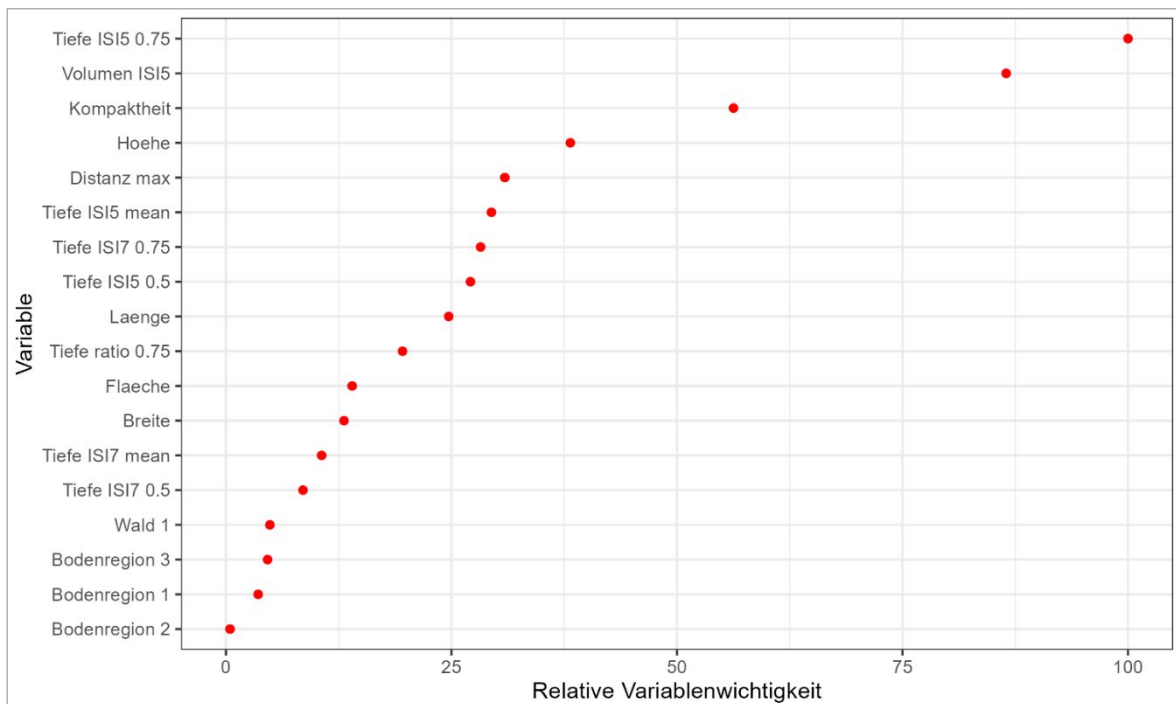


Abb. 14: Relative Variablenwichtigkeit des eingesetzten XGBoost-Modells für die Bodenregionen 1 bis 4.

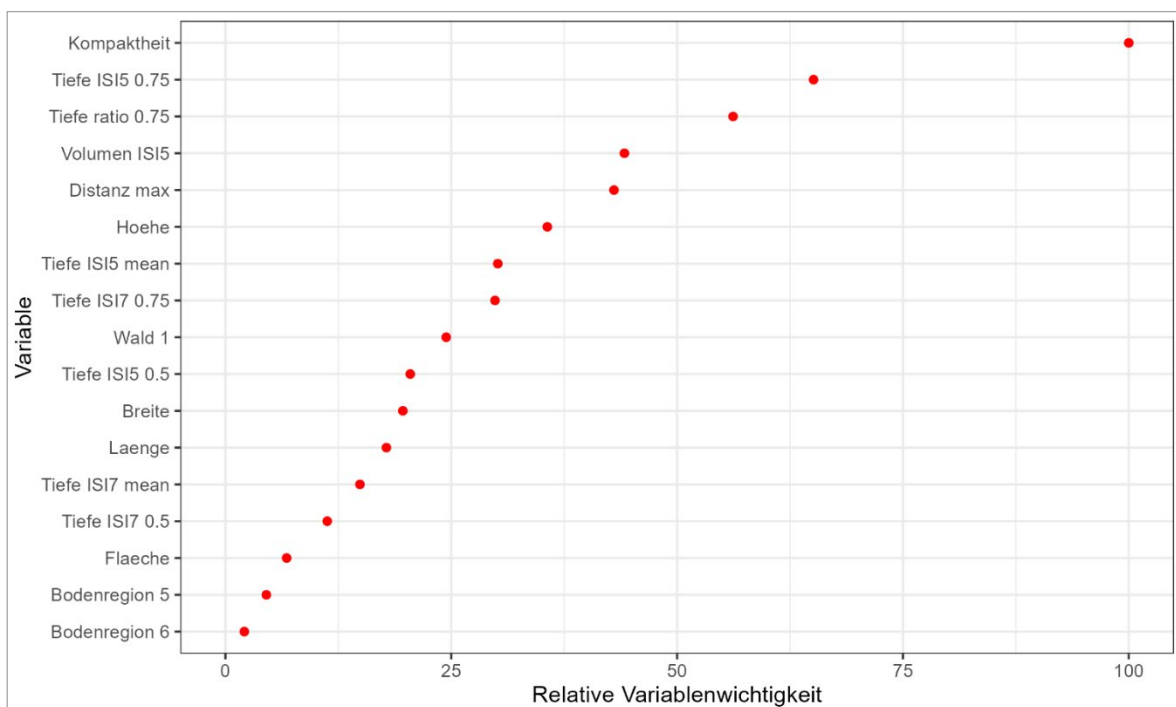


Abb. 15: Relative Variablenwichtigkeit des eingesetzten XGBoost-Modells für die Bodenregionen 5 und 6.

Eine weitere Möglichkeit, black-box-Modelle zu interpretieren, stellen die sogenannten SHapley Additive exPlanations (SHAP) dar, die von LUNDBERG & LEE (2017) entwickelt wurden und ihren Ursprung in der Spieltheorie haben (SHAPLEY 1953). Hierbei wird jeder Prädiktorvariablen (input feature) eines spezifischen Datenpunktes ein SHAP-Wert zugewiesen, welcher anzeigt, inwiefern die jeweilige Variable zur Vorhersage für diesen Datenpunkt beiträgt (MOLNAR 2023). SHAP ist demnach – im Gegensatz zur „klassischen“ Variablenwichtigkeit – nicht in erster Linie ein globaler Interpretationsansatz, sondern erlaubt lokale Aussagen für einzelne Datenpunkte, welche anschließend für eine globale Interpretation aggregiert werden können (MOLNAR 2022).

Die Abbildung 16 sowie die Abbildung 17 zeigen die Verteilungen der SHAP-Werte für die einzelnen Prädiktorvariablen innerhalb der XGBoost-Modelle. Für die Berechnungen waren zuvor verschiedene Methoden getestet worden, wobei sowohl modell-spezifische (LUNDBERG et al. 2019) als auch modell-agnostische (COVERT & LEE 2021) Varianten zum Einsatz kamen. Schließlich wurden die SHAP-Werte mit Hilfe des Paketes kernelshap (MAYER & WATSON 2024) gemäß COVERT & LEE (2021) berechnet und mit dem Paket shapviz (MAYER 2024) visualisiert.

Gemäß den beiden vorgenannten Abbildungen spielt die Kompaktheit der Cluster in beiden Modellen die größte Rolle bei der Vorhersage der drei möglichen Klassen „J“ (Graben), „N“ (kein Graben) und „U“ (unklar). Grundsätzlich wirken sich geringe Werte der Kompaktheit hierbei erwartungsgemäß positiv auf die Klasse „J“ und negativ auf die Klasse „N“ aus. Für die Klasse „U“, welche – wie bereits erwähnt – häufig flachere Gräben umfasst, ist zu erkennen, dass sich geringe Werte tendenziell positiv auswirken.

Ebenfalls in fast allen Fällen von hoher Relevanz ist die Prädiktorvariable „Tiefe ISI5 0.75“, wobei sich größere Werte tendenziell positiv auf die Vorhersage der Klassen „J“ und „U“ bzw. negativ auf die Klasse „N“ auswirken. Des Weiteren spielen die Variablen „Höhe“ und „Volumen“ zumeist eine wichtige Rolle. Tendenziell ist zu erkennen, dass sich eine größere Höhe negativ auf die Klassen „J“ und „U“ sowie positiv auf die Klasse „N“ auswirkt. Hierbei dürfte unter anderem die tendenziell mit der Höhe steigende Reliefenergie zum Tragen kommen. So treten bei

stark reliefierten Bereichen vergleichsweise viele Strukturen bzw. Cluster auf, die keine Gräben darstellen, sondern beispielsweise natürliche Geländeeinschnitte. Im Hinblick auf das Volumen der Cluster haben geringe Werte tendenziell eine negative Wirkung auf die Klasse „J“ und eine positive Wirkung auf die Klasse „N“. Dies kann im Einzelfall jedoch auch anders ausfallen, beispielsweise im Falle von tieferen bzw. größeren natürlichen Geländeeinschnitten.

Unterschiede zwischen den beiden Modellen bestehen beispielsweise hinsichtlich der Prädiktorvariablen „Tiefe ratio 0.75“, welche bei den Bodenregionen 1 bis 4 nur eine mittlere bzw. untergeordnete Rolle spielt, bei den Bodenregionen 5 und 6 hingegen eine deutlich höhere Bedeutung besitzt. Dies dürfte unter anderem darauf zurückzuführen sein, dass im Bergland bzw. im Harz deutlich mehr Hohlwege auftreten, bei denen die Variable tendenziell mittlere bis geringe Werte aufweist und welche vom Modell korrekterweise nicht als Gräben ausgewiesen werden.

Ein weiterer Unterschied zwischen den beiden Modellen ist etwa in Bezug auf die Prädiktorvariable „Wald“ zu erkennen. Zwar ist die „Wirkrichtung“ in beiden Modellen ähnlich, d. h. tendenziell wirkt sich die Lage im Wald negativ auf die Klasse „J“ und positiv auf die Klasse „N“ aus. Allerdings zeigt sich für die Bodenregionen 5 und 6 insgesamt eine deutlich höhere Bedeutung. Dies ist zum einen darauf zurückzuführen, dass diese beiden Bodenregionen insgesamt mehr Waldflächen aufweisen. Zum anderen dürfte sich hier – ähnlich zur Höhe – die Topographie widerspiegeln, wobei tendenziell stark reliefierte Bereiche häufiger mit Wald bestanden sind und viele Cluster aufweisen, die keine Gräben, sondern natürliche Geländeeinschnitte darstellen.

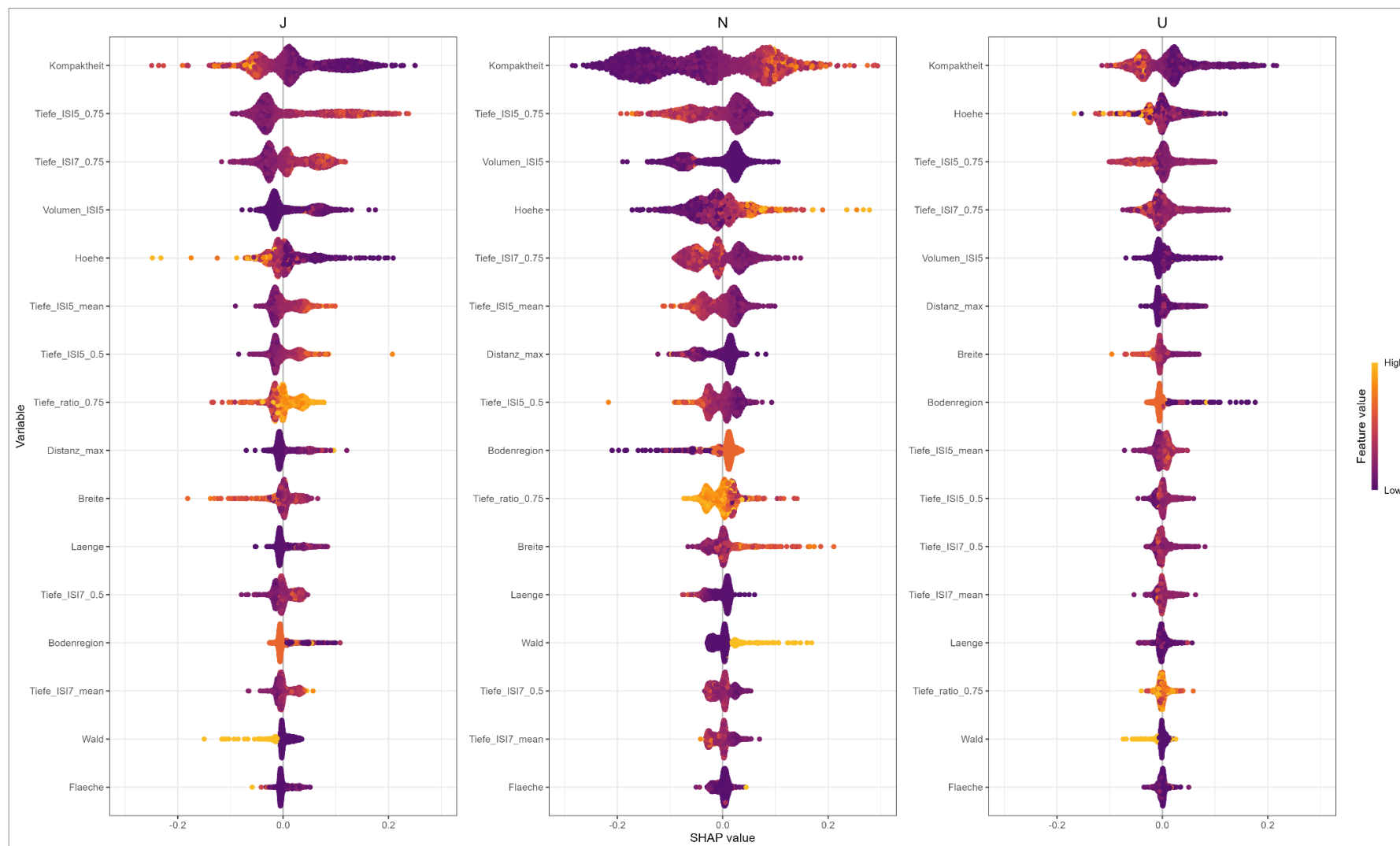


Abb. 16: SHAP-Werte der Prädiktorvariablen des XGBoost-Modells für die Bodenregionen 1 bis 4. Dargestellt sind die Verteilungen für die Vorhersageklassen J, N und U.

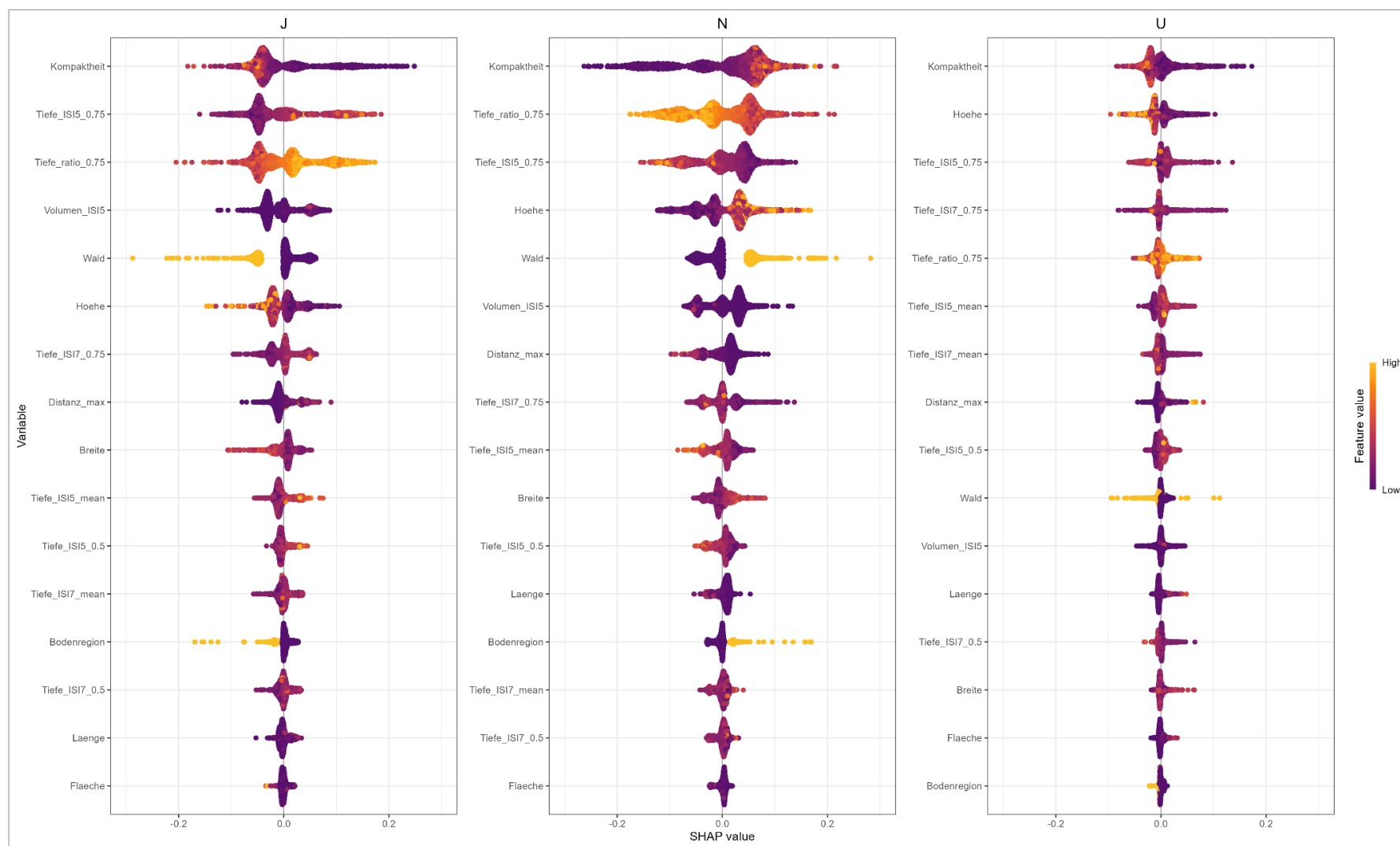


Abb. 17: SHAP-Werte der Prädiktorvariablen des XGBoost-Modells für die Bodenregionen 5 und 6. Dargestellt sind die Verteilungen für die Vorhersageklassen J, N und U.

6. Ergebnisse nach Durchführung des gesamten Workflows

6.1. Ergebnis nach Anwendung der XGBoost-Modelle

Nach dem Filtern der Ergebnisse wiesen die Rasterzellen eines jeden Clusters eine der drei Klassen „J“, „N“ und „U“ auf. Sämtliche weiteren Rasterzellen, die keinem Cluster zugehörig waren, erhielten ebenfalls die Zuordnung „N“. Die Abbildung 18 zeigt das Ergebnis für die beiden bereits erwähnten Ausschnitte aus den Trainingsgebieten 10 (Alte Kapelle) und 4 (Hüsedede).

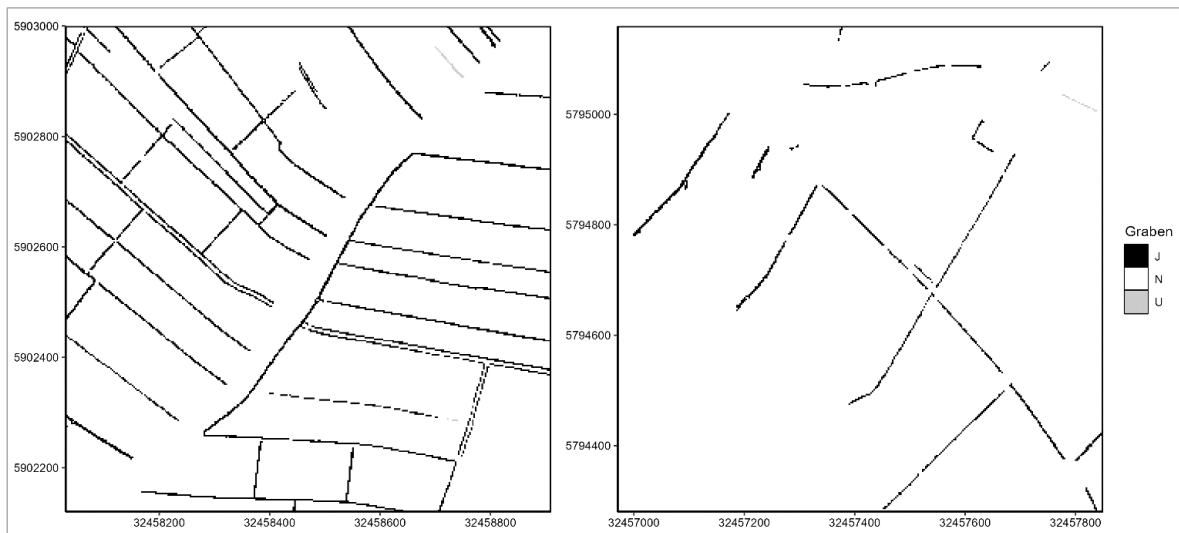


Abb. 18: Ergebnis nach dem Filtern. Links: Ausschnitt aus dem Trainingsgebiet 10 (Alte Kapelle), rechts: Ausschnitt aus dem Trainingsgebiet 4 (Hüsedede).

6.2. Abschätzung der Modellgüte für den gesamten Workflow

Um die Modellgüte des gesamten Workflows zu bewerten, sollen die Modellierungsergebnisse (nachfolgend auch als „prediction“ bezeichnet) mit den „tatsächlichen“ Daten („truth“) aus den Trainingsgebieten rasterzellenbasiert abgeglichen werden. Hierfür sind aus den nachfolgend genannten Gründen jedoch zunächst einige Bearbeitungsschritte erforderlich bzw. sinnvoll. Zum einen weisen die Rasterzellen im Ergebnis

eine der drei Klassen „J“, „N“ und „U“ auf, während im Zuge der Digitalisierung bzw. des Modelltrainings lediglich die beiden Klassen „1“ und „0“ verwendet wurden. Zum anderen hatten bei der Erstellung der Trainingsdaten nur die Rasterzellen in unmittelbarer Nähe der Grabensohle das Label „1“ (Graben) erhalten, die Böschungsbereiche blieben unberücksichtigt (s. Abschnitt 3.3). In der Modellvorhersage wurden jedoch auch etwaige Böschungsbereiche berücksichtigt bzw. beibehalten, um Informationen über die tatsächliche Ausdehnung der modellierten Gräben an jeder Stelle zu erhalten. Bei

einem direkten Vergleich würde die Anzahl an „falsch positiven“ Rasterzellen (FP) daher vermutlich stark überschätzt. Die Abbildung 19 verdeutlicht die Unterschiede exemplarisch an einem Ausschnitt aus dem Trainingsgebiet 3 (Venner Mühlenbach).

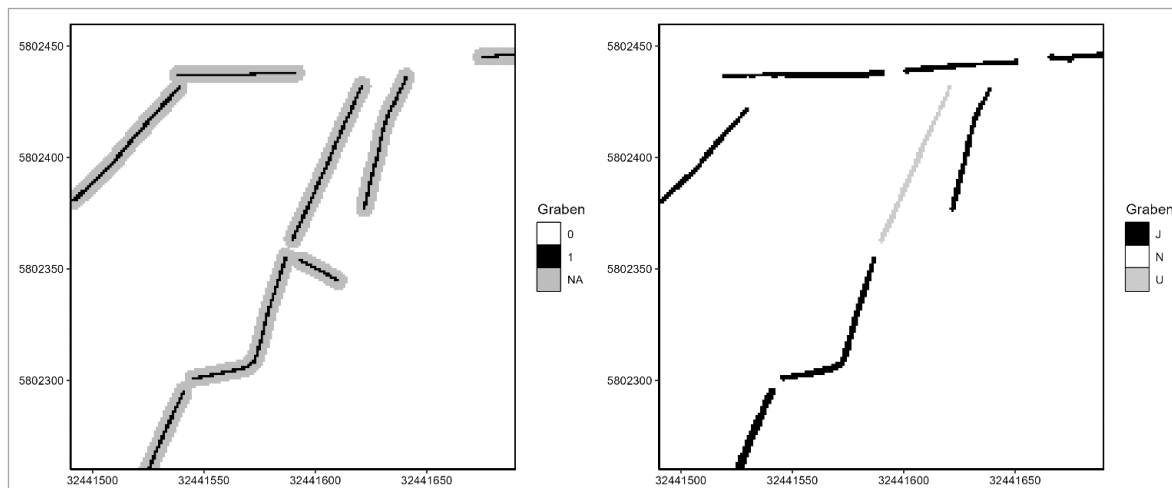


Abb. 19: Klassenzugehörigkeit der Rasterzellen. Links: Tatsächlich digitalisierte Daten („truth“), rechts: Modellergebnis („prediction“).

Es wurden daher die folgenden Schritte nacheinander durchgeführt.

- **Anpassung der Ergebnisklassen:**
Dieser Schritt erfolgte, um die tatsächlichen Daten mit den Modellergebnissen vergleichen zu können. Zunächst erhielten die Rasterzellen, welchen im Ergebnis die Klasse „J“ zugewiesen worden war, die Klasse „1“. Analog wurden für die Berechnung auch die Rasterzellen der Klasse „U“ der Klasse „1“ zugeordnet, da es sich bei diesen Strukturen augenscheinlich überwiegend um flache Gräben handelt, die vom Modell erkannt wurden.
- **Zuschnitt der modellierten Gräben auf die Zentrallinien:**
Durch diesen Schritt sollte insbesondere eine starke Überschätzung falsch positiver Rasterzellen vermieden werden. Um die

modellierten Gräben, welche auch bei der Digitalisierung erfasst worden waren, auf ihre Zentrallinie zuzuschneiden, wurden die Rasterzellen, welche ursprünglich nicht berücksichtigt worden waren (NA), im Ergebnistraster zunächst mit dem Label „0“ versehen. Anschließend wurden die verbleibenden Rasterzellencluster rechnerisch mit Hilfe des Paketes whitebox (WU & BROWN 2022) auf eine Zentrallinie reduziert. Von den ggf. zusätzlich modellierten Strukturen, welche im Rahmen der Digitalisierung nicht als Graben definiert worden waren, wurde ebenfalls eine rechnerische Zentrallinie erzeugt.

Die nachfolgende Abbildung 20 zeigt das Ergebnis nach der Durchführung der beiden vorgenannten Schritte für den bereits erwähnten Ausschnitt aus dem Trainingsgebiet 3.

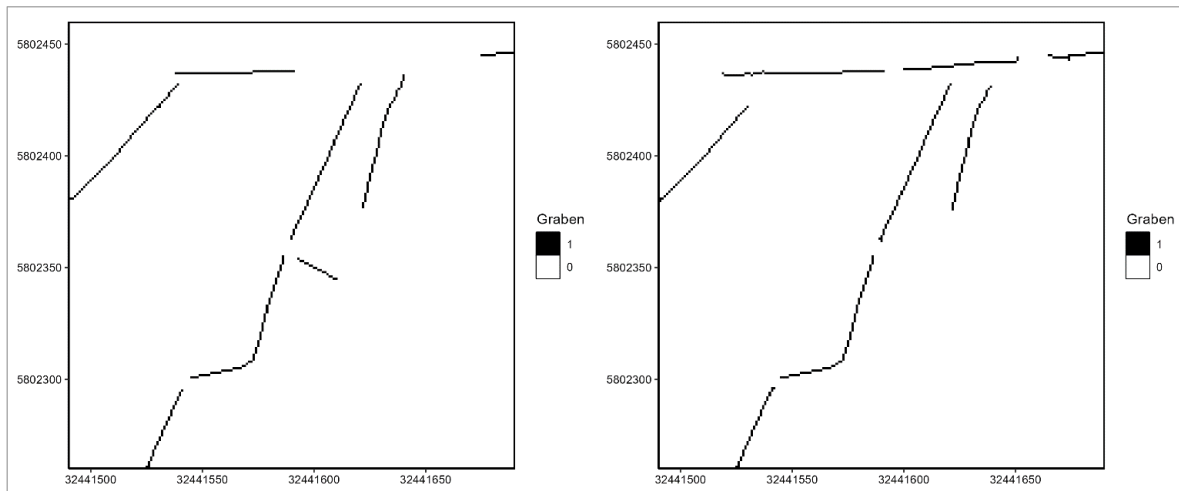


Abb. 20: Vergleich der tatsächlichen Daten („truth“, links) mit dem Modellergebnis („prediction“, rechts) nach der Anpassung der Klassen und des Zuschnittes der Grabencluster auf ihre Zentrallinien.

Die nunmehr möglichen vier Kombinationen aus tatsächlichen und modellierten Daten sind in der Abbildung 21 dargestellt. Hierbei stellen Rasterzellen mit dem Label „TP“ (true positive) korrekt vorhergesagte „Grabenpixel“ dar, während das Label „TN“ (true negative) die korrekt modellierten „Nicht-Graben-Pixel“ beschreibt. Demgegenüber wurden tatsächliche „Grabenpixel“, welche vom Modell nicht als solche erkannt wurden, als „FN“ (false negative) bezeichnet. Wurden Rasterzellen als „Grabenpixel“ ausgewiesen, die zuvor nicht digitalisiert worden waren, erhielten sie das Label „FP“ (false positive).

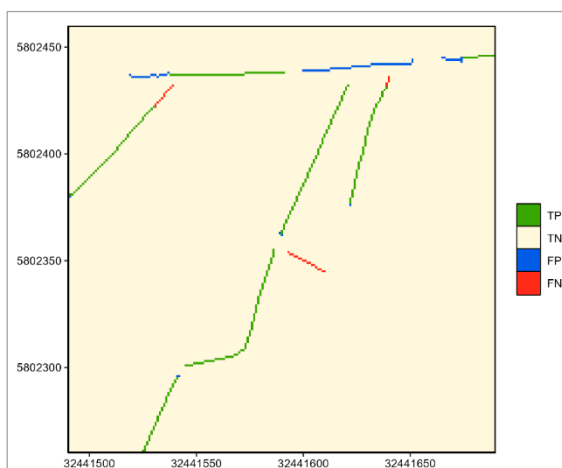


Abb. 21: Die vier möglichen Kombinationen aus tatsächlichen und modellierten Daten werden durch die Label „TP“ (true positive), „TN“ (true negative), „FP“ (false positive) sowie „FN“ (false negative) dargestellt.

Für die Bewertung der Modellgüte bei einer binären Klassifikation stehen verschiedene Evaluationsmetriken zur Verfügung. Im vorliegenden Fall wurden die drei folgenden Metriken ausgewählt und mit dem Paket *cvms* (OLSEN & ZACHARIAE 2024) berechnet:

■ Recall (auch: Sensitivity):

Dieser Parameter gibt an, wieviele der tatsächlichen „Grabenpixel“ auch als solche vom Modell erkannt wurden und stellt somit ein Maß für die Vollständigkeit der Vorhersage dar:

$$Recall = \frac{TP}{TP + FN}$$

■ Precision:

Hierbei werden die korrekt erkannten „Grabenpixel“ ins Verhältnis zu allen vorhergesagten „Grabenpixeln“ gesetzt, so dass sich ein Maß für die Genauigkeit der Vorhersage ergibt:

$$Precision = \frac{TP}{TP + FP}$$

■ Kappa-Index (auch: Cohen's Kappa):

Dieser Index gibt an, inwieweit das Ergebnis von einer rein zufälligen Verteilung abweicht (s. Anmerkungen im Abschnitt 3.4). Nach einer Normalisierung von TP, TN, FP und FN auf eine Gesamtsumme von 1 ermittelt sich der Kappa-Index folgendermaßen (OLSEN & ZACHARIAE 2024):

$$P_{observed} = TP + TN$$

$$P_{expected} = (TN + FP)(TN + FN) + (FN + TP)(FP + TP)$$

$$Kappa = \frac{P_{observed} - P_{expected}}{1 - P_{expected}}$$

Die Abbildung 22 zeigt die Verteilung der tatsächlichen und vorhergesagten Daten für alle Rasterzellen der insgesamt neun betrachteten Trainingsgebiete in einer sogenannten „confusion matrix“. Des Weiteren sind die drei vorgenannten Evaluationsmetriken angegeben. Hierbei zeigt sich, dass der überwiegende Anteil der insgesamt 36.000.000 Rasterzellen korrekterweise keinem Graben zugeordnet wurde

(99,45 % TN). Zudem wurden die digitalisierten Grabenstrukturen überwiegend vom Modell erkannt und nur in untergeordnetem Maße nicht ausgewiesen; die Wiederfindungsrate („recall“) ergibt sich zu rund 88 %. Teilweise wurden auch Grabenstrukturen modelliert, die bei der Digitalisierung nicht berücksichtigt worden waren. Dies ist allerdings nur in vergleichsweise geringem Umfang der Fall und äußert sich in einer Genauigkeit („precision“) von rund 83 %. Auch der Kappa-Index von 0,85 weist auf eine sehr hohe Übereinstimmung von Realität und Modellergebnis hin. Im Vergleich zu anderen Studien zur automatisierten Grabenerkennung fallen die Werte ebenfalls sehr gut aus (s. Übersicht in LIDBERG et al. (2023)).

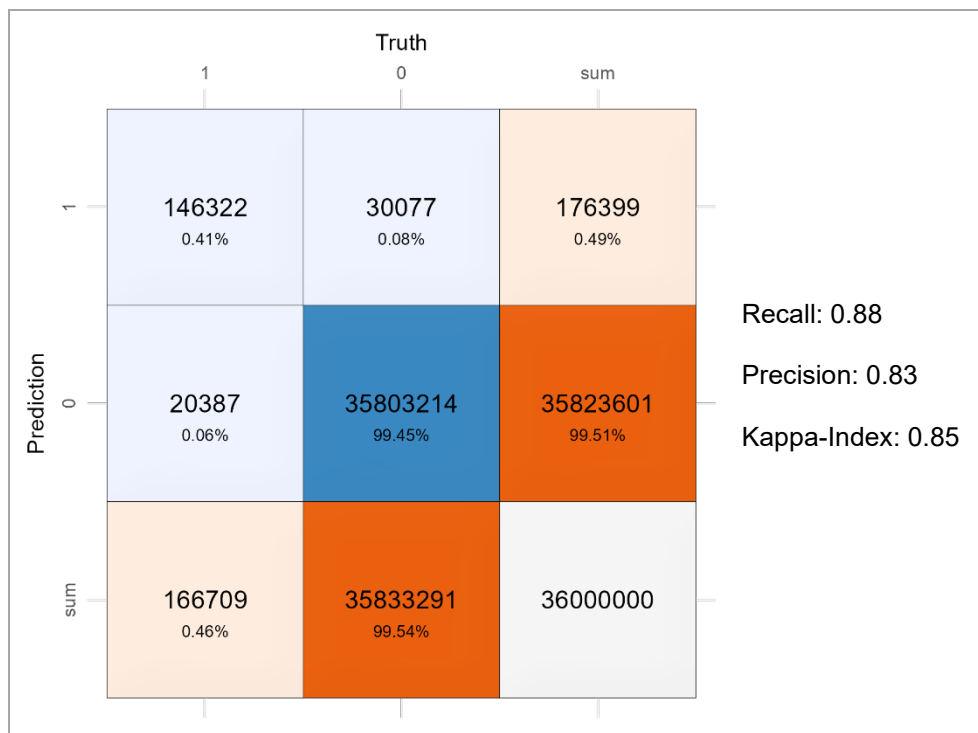


Abb. 22: Confusion matrix und Evaluationsmetriken für sämtliche Rasterzellen der neun betrachteten Trainingsgebiete.

Werden die einzelnen Trainingsgebiete separat betrachtet, ergeben sich die in der Tabelle 6 genannten Metriken. Die Werte liegen hierbei insgesamt auf einem hohen bis sehr hohen Niveau. Teilweise treten gewisse Abweichungen in den Trainingsgebieten 2 und 3 sowie 6 und 7 auf. Im Falle der Trainingsgebiete 2 und 3 liegt

der precision-Wert bei 0,67 bzw. 0,70 und damit niedriger als in den anderen Gebieten. Dies ist darauf zurückzuführen, dass neben den tatsächlich digitalisierten Gräben vergleichsweise viele weitere Strukturen vom Modell als Graben erkannt wurden. Dies ist nach augenscheinlicher Überprüfung jedoch kein Problem, da es

sich bei diesen Strukturen größtenteils tatsächlich um Entwässerungsgräben handelt. In den Trainingsgebieten 6 und 7 wurden mit 0,51 bzw. 0,72 vergleichsweise geringe recall-Werte ermittelt. Diese ergeben sich, da vor allem in den mit Wald bestandenen Gebieten einige der zunächst digitalisierten Gräben nicht vom Modell erkannt wurden. Dies ist nicht optimal, jedoch

insgesamt als nicht kritisch zu bewerten, da es sich hierbei häufig um vergleichsweise flache Gräben, häufig am Rand von Straßen bzw. Wegen handelt. Generell bestätigt dies die Erwartung, dass die korrekte Grabenerkennung in bewaldeten Bereichen tendenziell schwieriger ist als im Offenland.

Tab. 6: Evaluationsmetriken für die einzelnen Trainingsgebiete. Nähere Informationen zu den jeweiligen Gebietseigenschaften gehen aus der Tabelle 2 hervor.

Trainingsgebiet		Recall	Precision	Kappa-Index
Nr.	Bezeichnung			
1	Wimmerbach	0.95	0.88	0.92
2	Kronensee	0.96	0.67	0.79
3	Venner Mühlenbach	0.90	0.70	0.78
4	Hüsedde	0.94	0.76	0.84
5	Schickelsheim	0.86	0.94	0.90
6	Schieringen	0.51	0.92	0.66
7	Raebke	0.72	0.89	0.80
8	Aschhauserfeld	0.79	0.81	0.80
10	Alte Kapelle	0.99	0.90	0.94

6.3. Verteilung der modellierten Strukturen in der Fläche

Den einzelnen Rasterzellen war im Verlauf der Modellierung eine der drei Klassen „J“ (Graben), „N“ (kein Graben) bzw. „U“ (unklar, Grenzfall) zugewiesen worden. Nachfolgend werden die Klasse „J“ auch als „Grabenstruktur wahrscheinlich“ und die Klasse „U“ als „Grabenstruktur möglich“ interpretiert. Wenn die modellierten „Grabenpixel“ ins Verhältnis zur Gesamtzahl an Rasterzellen gesetzt werden, ergibt sich eine

rechnerische Grabendichte. Über die gesamte Landesfläche errechnet sich ein Wert von etwa 1,2 %, wobei hiervon rund 1,1 Prozentpunkte auf die Einstufung „Grabenstruktur wahrscheinlich“ entfallen. Die Rasterzellen, welche als „Grabenstruktur möglich“ einzustufen sind, machen einen lediglich geringen Anteil von < 0,1 % aus. Die Abbildung 23 zeigt die entsprechenden Werte für die einzelnen Bodenregionen. Eine Übersichtskarte der Bodenregionen ist ergänzend als Abbildung 24 dargestellt.

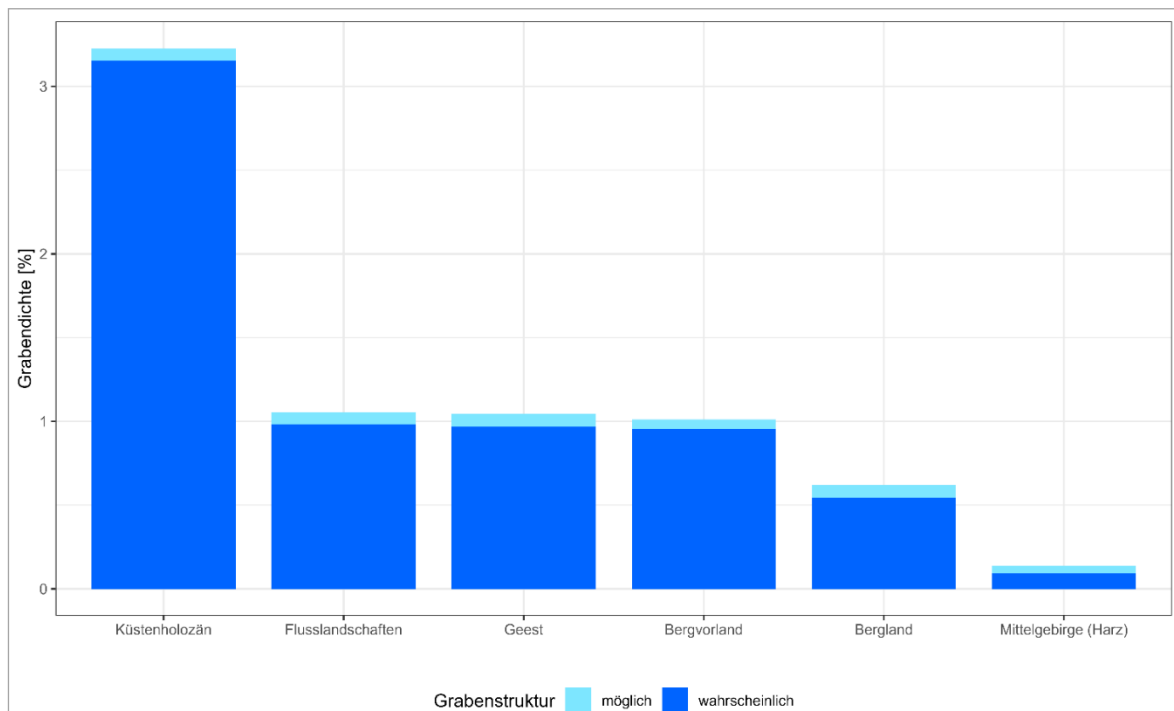


Abb. 23: Grabendichte der sechs Bodenregionen in %.

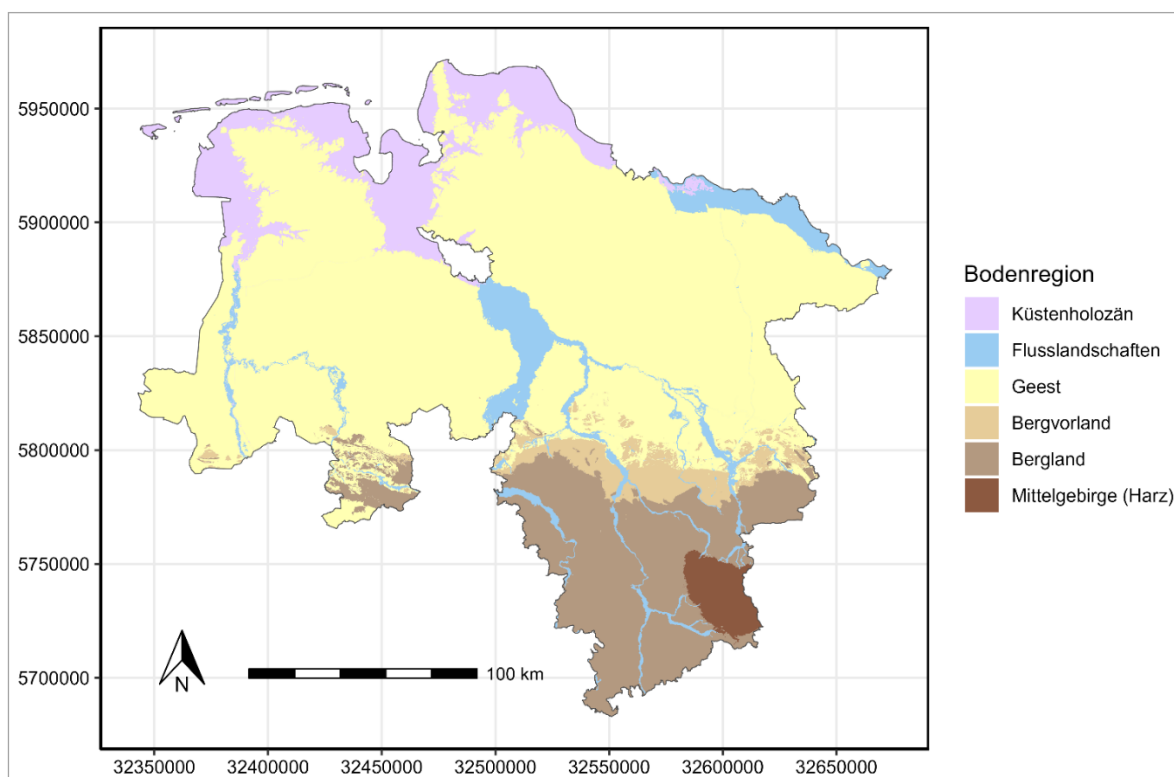


Abb. 24: Bodenregionen in Niedersachsen (GEHRT et al. (2021), verändert).

Demnach weist das Küstenholozän die mit deutlichem Abstand höchste Grabendichte auf. Rund 3,2 % der Rasterzellen werden hier als „Grabenstruktur wahrscheinlich“ ausgewiesen. Im Bereich der Flusslandschaften sowie der Geest und des Bergvorlandes ergeben sich die entsprechenden Grabendichten zu jeweils rund 1,0 %, für das Bergland zu rund 0,5 % und für den Bereich Harz zu rund 0,1 %. Der Anteil an Rasterzellen, welche als „Grabenstruktur möglich“ ausgewiesen wird, liegt für die verschiedenen Bodenregionen zwischen rund 0,04 % und 0,08 %.

Werden die Daten für die betrachteten 1 x 1 km großen Kacheln aggregiert, kann die Verteilung innerhalb der Bodenregionen besser abgebildet werden. Die nachfolgende Abbildung 25 zeigt entsprechende Box-Whisker-Plots für die Rasterzellen, bei denen eine Grabenstruktur wahrscheinlich ist. In den Plots stellen die untere und die obere Begrenzung der Box das untere bzw. das obere Quartil (0,25-Quantil bzw. 0,75-Quantil) der Daten dar. Der Median ist als waagerechter Strich dargestellt. Für eine bessere Übersichtlichkeit wurden nur die Datenpunkte visualisiert, welche zwischen dem 0,05- und dem 0,95-Quantil liegen. Diese werden durch die Enden der beiden Whisker (Antennen) symbolisiert.

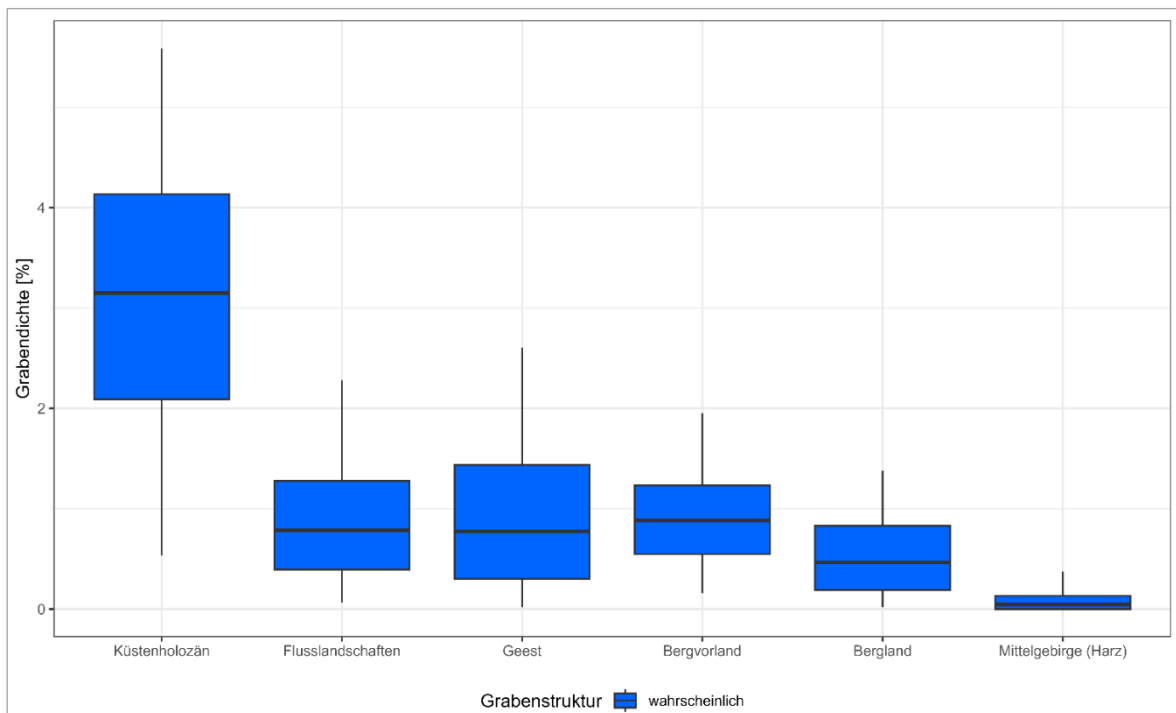


Abb. 25: Grabendichte der sechs Bodenregionen, aggregiert für die betrachteten 1 x 1 km großen Kacheln.

Die Abbildung 26 zeigt die Verteilung dieser Grabendichte für die Landesfläche. Auch für diese Abbildung wurden nur die Strukturen mit der Bewertung „Grabenstruktur wahrscheinlich“ berücksichtigt.

Hierbei treten die räumlichen Unterschiede der Grabendichte insgesamt gut hervor. So ist vor allem der Bereich des Küstenholozäns gut an einer hohen Grabendichte zu erkennen. Die Geestbereiche fallen durch eine vergleichsweise hohe Bandbreite hinsichtlich der Grabendichte auf. Zum einen sind beispielsweise im nordwestlichen Bereich tendenziell höhere Grabendichten als in den übrigen Geestgebieten zu verzeichnen. Zum anderen zeigt sich, dass im

östlichen Teil des Landes nur vergleichsweise wenige Gräben auftreten. In diesen Ausprägungen spielen vermutlich sowohl bodenkundliche als auch klimatische Gegebenheiten eine Rolle. Zudem sind lokal einige sehr hohe Grabendichten zu erkennen, vor allem im mittleren Landes-
teil. Hierbei handelt es sich häufig um (ehemalig) genutzte Moore, welche intensiv entwässert werden bzw. wurden. Dies kann auch in aufgeforsteten Bereichen auftreten. Im Bergland ist insgesamt eine etwas geringere Grabendichte zu verzeichnen, wobei sich die einzelnen Höhenzüge häufig mit deutlich reduzierten Werten abheben. Ebenfalls fällt der Harz durch besonders geringe Grabendichten auf.

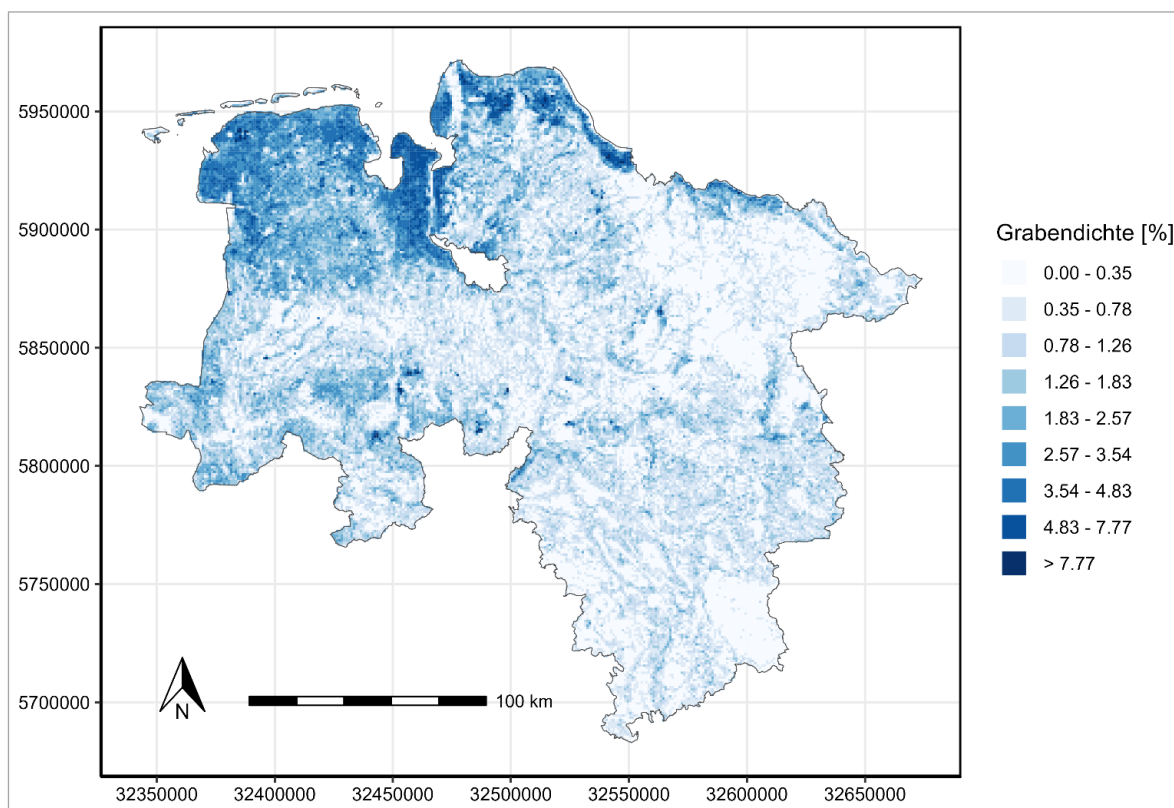


Abb. 26: Grabendichte für die im Modell verwendeten Kacheln von jeweils 1 x 1 km Größe.

7. Diskussion und Ausblick

7.1. Zusammenfassung

Bei der (halb-)automatisierten Erkennung von Entwässerungsgräben kommen vermehrt Methoden des Maschinellen Lernens (ML) zum Einsatz (BALADO et al. 2019; DU et al. 2024; FLYCKT et al. 2022; LIDBERG et al. 2023; ROBB et al. 2023; ROELEN, ROSIER et al. 2018). Die hier vorgestellte Methode kombiniert mit Random Forest sowie XGBoost zwei dieser Ansätze für die Gebietskulisse des Landes Niedersachsen, wobei insgesamt gute bis sehr gute Ergebnisse erzielt werden konnten. Die Evaluationsmetriken fallen auch im Vergleich mit anderen Studien (s. z. B. Zusammenstellung in LIDBERG et al. (2023)) sehr gut aus. Dabei umfasst das Anwendungsgebiet mit rund 47.700 km² nicht nur ein Vielfaches der in anderen Studien betrachteten Flächen, sondern zeichnet sich zudem durch eine ausgeprägte landschaftliche Heterogenität aus.

Die hauptsächliche Herausforderung bei der Modellierung bestand darin, zum einen möglichst viele Gräben mit unterschiedlichen Längen, Breiten, Tiefen, Querprofilen und Verläufen zu erkennen, und zum anderen möglichst viele andere Strukturen trotz möglicherweise ähnlicher Form oder Lage auszuschließen. Die Anwendung eines auf verschiedenen Terrainindices beruhenden Random-Forest-Modells lieferte hierbei eine Vorhersage, welche für jede Rasterzelle eine „Grabenwahrscheinlichkeit“ ausweist. In ebenen und flach geneigten landwirtschaftlich genutzten Bereichen stimmten die

Rasterzellen mit einer hohen Grabenwahrscheinlichkeit gut mit den tatsächlich vorhandenen Gräben überein. In Gebieten mit erhöhter Reliefenergie kamen die vorhandenen Gräben in der Regel zwar ebenfalls zur Geltung; allerdings wurden – vor allem im Bereich des Berg- und Hügellandes sowie in bewaldeten Gebieten – vergleichsweise viele Rasterzellen mit hohen Grabenwahrscheinlichkeiten ausgewiesen, welche tatsächlich zu Hohlwegen, natürlichen Geländeeinschnitten und Gewässern gehörten. Dieser Problematik wurde versucht mit verschiedenen Post-Processing-Schritten sowie dem Einsatz eines auf XGBoost basierenden Filters zu begegnen, was zu einer deutlichen Verbesserung der Vorhersage führte. Die Abbildung 27 zeigt Beispiele für Strukturen, die mit dem entwickelten Workflow korrekt erkannt wurden. Demnach wurden Gräben mit unterschiedlichen Breiten, Tiefen und Querprofilen sowie in verschiedenen Landschafts- und Nutzungseinheiten erfasst.

Da die Modellierung im Wesentlichen auf dem DGM1 basiert, werden bei der Modellierung nur Strukturen erfasst, die sich an der Geländeoberfläche zeigen. Unterirdische Gewässerabschnitte bei Durchlässen, Überfahrten, Verrohrungen o. ä. werden durch das Modell nicht erkannt. Das Ergebnis ist – analog zum DGM1 – ein Rasterdatensatz mit einer Rasterzellengröße von 1 x 1 m. Es werden zusammenhängende Cluster von Rasterzellen ausgewiesen, welche entweder mit dem Label „Grabenstruktur wahrscheinlich“ oder „Grabenstruktur möglich“ gekennzeichnet sind.

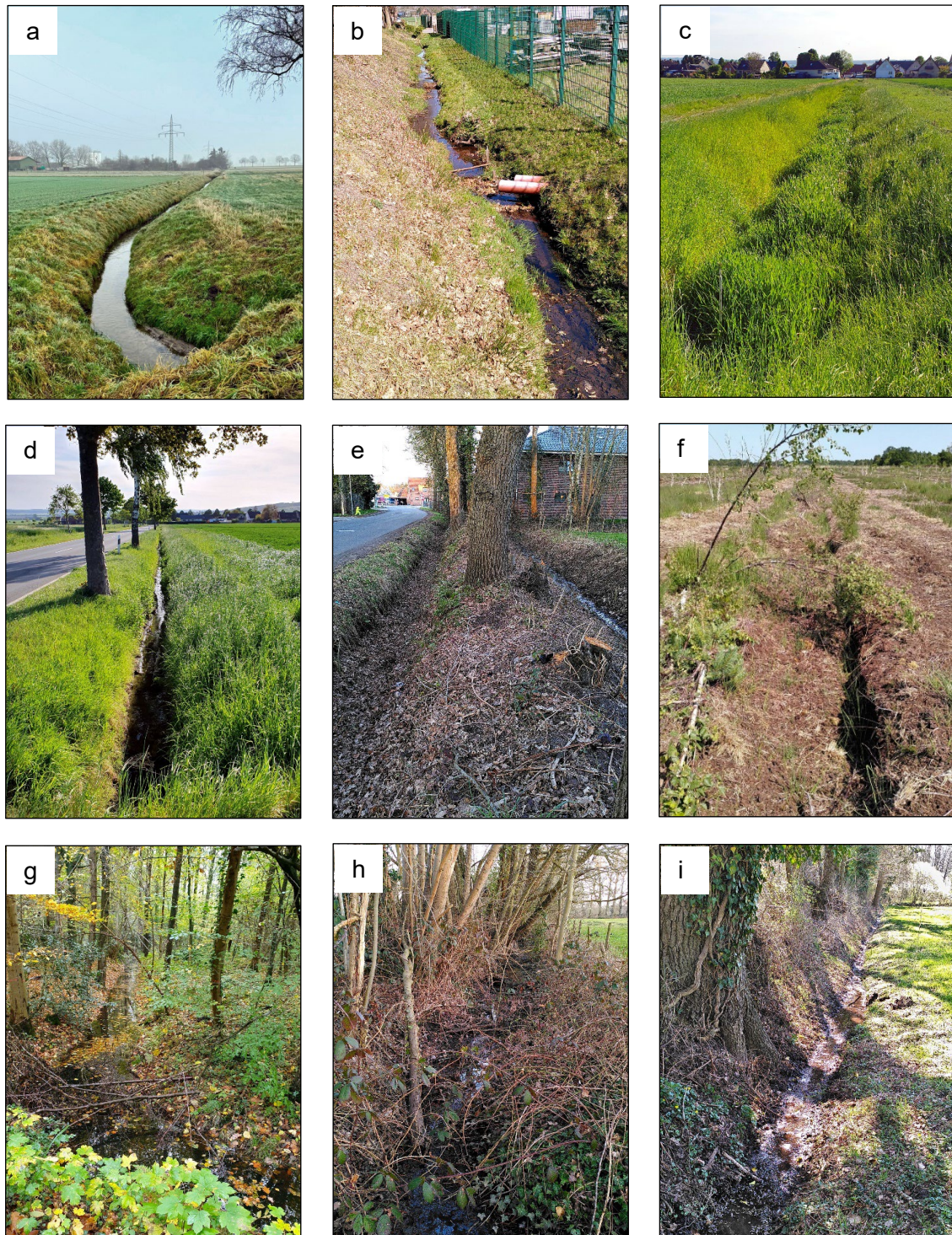


Abb. 27: Beispiele für Gräben, die mit der angewandten Methode korrekt erkannt wurden (a bis h: Grabenstruktur wahrscheinlich, i: Grabenstruktur möglich). Fotos: G. Griffel (a), J. Wessels (b, e, f, h, i), M. Hoetmer (c, d, g).

7.2. Möglichkeiten und Grenzen

Mit dem im vorliegenden Bericht vorgestellten Ansatz können Entwässerungsgräben anhand des DGM1 sowie weiterer Daten (halb)automatisiert erkannt werden. Mit dem rasterzellenbasierten Ergebnis kann neben der Erkennung bislang nicht kartierter Gräben auch die Lage bekannter Strukturen – inklusive einer ungefähren Länge und Breite – gut nachvollzogen werden. Die entwickelte Methode stellt einen deutlichen Vorteil gegenüber einer flächenhaften händischen Kartierung vor Ort dar, welche weit aus personal- und kostenintensiver ausfallen dürfte. Ein weiterer Vorteil gegenüber klassischer Kartierung ist die Möglichkeit, etwaige Veränderungen in der Landschaft automatisiert erfassen zu können. Voraussetzung hierfür ist eine möglichst hohe Aktualität bzw. Aktualisierungsrate der Eingangsdaten, insbesondere des Digitalen Geländemodells. Nach Angaben des Landesamtes für Geoinformation und Landesvermessung Niedersachsen (LGLN) sollen die zur DGM-Erstellung benötigten Daten des Airborne Laserscannings (ALS) ab 2025 in einem sechsjährigen Turnus erhoben werden (HÖNNIGER 2024). Dies stellt eine erhebliche Verbesserung zum bisherigen Stand dar und ermöglicht somit perspektivisch eine schnellere Aktualisierung auch bei der Erkennung von Entwässerungsgräben.

Generell ist hierbei jedoch zu berücksichtigen, dass es sich bei dem Verfahren um ein Modell handelt, welches nicht fehlerfrei ist. Insbesondere aufgrund der Heterogenität der niedersächsischen Landschaft und der vielfältigen Ausprägung der vorhandenen Entwässerungsgräben bleiben gewisse Herausforderungen und Fehlerquellen bestehen. Im Wesentlichen können hierbei zwei Varianten unterschieden werden:

Variante 1: Vorhandene Gräben werden nicht erkannt

Nicht alle tatsächlich vorhandenen Gräben konnten mit dem entwickelten Verfahren erkannt werden. Dies ist zum einen grundsätzlich darauf zurückzuführen, dass die Zuordnung einer Struktur zu einem Entwässerungsgraben selbst bei einer Prüfung im Gelände nicht immer eindeutig möglich ist. So treten landesweit praktisch sämtliche Variationen sowie Übergangsformen zwischen Gräben, Entwässerungsmulden, natürlichen Gewässern und Kanälen auf. Im Ergebnis bedeutet dies im vorliegenden Fall, dass vor allem vergleichsweise kleine, flache Strukturen, die beispielsweise bei Starkregenereignissen Oberflächenabfluss aufnehmen bzw. abführen können, allerdings im Normalfall keine nennenswerte Drainagewirkung aufweisen, nicht als Gräben erkannt werden (s. Abschnitt 5.2 bzw. Abbildung 13). Auch sehr breite Strukturen mit einer großen freien Wasserfläche wurden in der Regel nicht mehr als Gräben ausgewiesen, da sie häufig nicht das eingangs erwähnte typische Querprofil eines Grabens aufweisen. In beiden vorgenannten Fällen ist eine Nicht-Erkennung also in gewisser Weise beabsichtigt bzw. folgerichtig. Zudem werden vergleichsweise sehr kurze Grabenabschnitte zum Teil im Zuge des Post-Processings herausgefiltert. Dies kann beispielsweise bei vielen Unterbrechungen durch Überfahrten oder Verrohrungen auftreten.

Zum anderen ist die Modellierung stets so aktuell wie die Eingangsdaten. Für einzelne Bereiche der Landesfläche reichen die Aufnahmezeiträume des Laserscannings für die DGM-Erstellung beispielsweise bis in das Jahr 2013 zurück. Es ist davon auszugehen, dass sich – zumindest bereichsweise – zwischenzeitlich gewisse Änderungen als Folge von Pflege- und/oder Vertiefungsmaßnahmen an bestehenden Gräben bzw. durch die Anlage neuer Gräben ergeben haben. Hier sind etwa Maßnahmen im Anschluss an das niederschlagsreiche Winterhalbjahr 2023/2024 zu nennen.

Die Abbildung 28 zeigt Strukturen, welche bei der Modellierung nicht als Graben erkannt wurden.



Abb. 28: Beispiele für Strukturen, die nicht als Graben erkannt wurden. a: Grabenabschnitte sind vergleichsweise kurz, da viele Überfahrten vorhanden sind, b: Struktur ist sehr breit und weist kein typisches Grabenprofil auf, c: Graben kam im DGM1 kaum heraus, vermutlich haben zwischenzeitlich Vertiefungsmaßnahmen stattgefunden. Fotos: C. Schimschal (a), J. Wessels (b), G. Ertl (c).

Variante 2: Andere Strukturen werden als Gräben ausgewiesen

Trotz eines intensiven Post-Processings wurden bei der Modellierung teilweise auch andere in der Landschaft vorkommende Elemente als Gräben ausgewiesen. Dies betrifft beispielsweise schmale natürliche Gewässer sowie anthropogene Strukturen wie Hohlwege und in

das Gelände eingeschnittene schmale Straßen und Einfahrten. Als Ursachen hierfür kommen im Wesentlichen die längliche Form der Elemente sowie die insgesamt vergleichbaren räumlichen Ausdehnungen und/oder Querprofile in Betracht. Die Abbildung 29 zeigt Beispiele für Strukturen, die vom Modell fälschlicherweise als Gräben erkannt wurden.

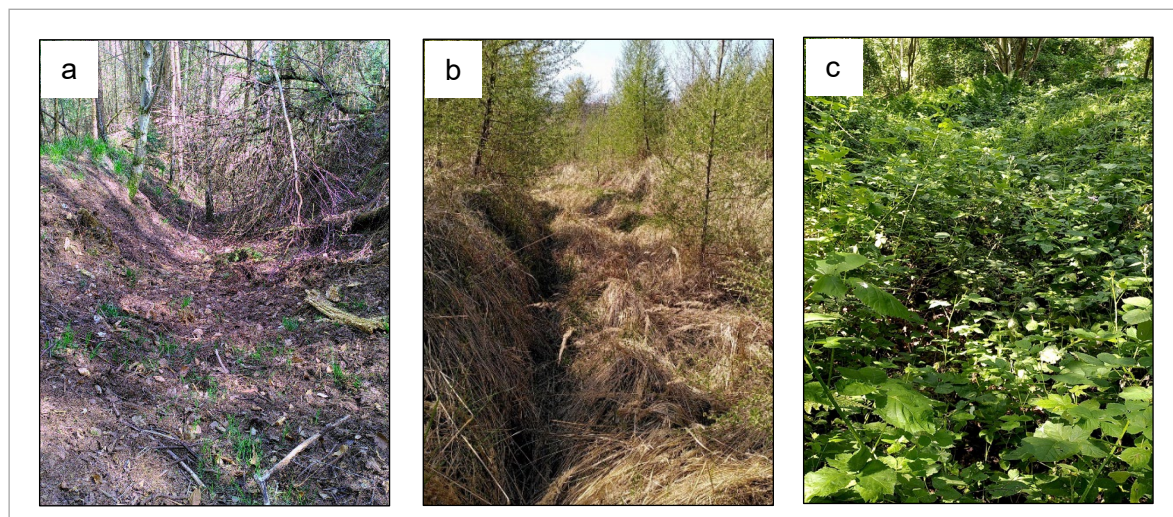


Abb. 29: Strukturen, die fälschlicherweise als Graben erkannt wurden. a: Hohlweg, b: eingeschnittener ehemaliger Weg mit tiefen Fahrspuren, c: Geländeeinschnitt in ehemaliger Sandgrube. Fotos: J. Wessels.

Insgesamt werden die vorgenannten Einschränkungen vor dem Hintergrund der Größe des Untersuchungsgebietes, der heterogenen topographischen Gegebenheiten in Niedersachsen sowie der zur Verfügung stehenden Mittel derzeit als vertretbar eingestuft. Das Ziel der vorliegenden Arbeit, eine plausible Hinweiskarte für Entwässerungsgräben zu erstellen, wurde erreicht.

Im Hinblick auf die technischen Rahmenbedingungen bei der Anwendung von Methoden des Maschinellen Lernens sei an dieser Stelle darauf hingewiesen, dass es sich bei dem hier vorgestellten Workflow um einen „hands-on“-Ansatz handelt, welcher in erster Linie im Hinblick auf seine Funktionalität entwickelt wurde. So ist z. B. zu berücksichtigen, dass sowohl bei dem eingesetzten Random-Forest-Modell als auch bei den im Post-Processing entwickelten XGBoost-Modellen jeweils Variablen mit unterschiedlichen Varianten eingesetzt wurden. Es lag somit jeweils – zumindest teilweise – Multikollinearität vor, was etwa bei der Interpretierbarkeit der SHAP-Werte ggf. zu gewissen Einschränkungen führen kann (MOLNAR 2023). Diesem Punkt wurde u. a. durch die randomisierte Auswahl einer begrenzten Anzahl an Prädiktorvariablen beim Ausbau der Entscheidungs- bzw. Klassifikationsbäume zu begegnen versucht. Die verbliebene „Unsicherheit“ wurde in Kauf genommen, weil sich das Modellergebnis unter Bezugnahme der jeweiligen Variablen bzw. Varianten verbesserte.

Die entwickelte Methode ist prinzipiell auf andere Gebiete mit vergleichbarer Datenbasis übertragbar. Zudem kann eine vergleichsweise schnelle Anwendung auf große Gebiete erfolgen. Nach Einschätzung der Autoren ist hierbei eine Anpassung auf die jeweiligen naturräumlichen Gegebenheiten erforderlich bzw. sinnvoll.

Mit den Modellergebnissen können wichtige Informationen zur Entwässerung der Landschaft gewonnen werden. Die Grabendichte in Abbildung 26 gibt hierzu einen ersten Hinweis. Weitere Auswertungen sind u. a. im Rahmen des Projektes KliBoG1 geplant. Zudem kann die Kenntnis über vorhandene Grabenstrukturen auch für andere Bereiche wichtige Informationen liefern. Zu nennen sind hier beispielsweise Möglichkeiten zum Wasserrückhalt in der Fläche, etwa über den Einstau von Gräben zur Renaturierung von Mooren (GRAF et al. 2022). Des Weiteren können Gräben eine wichtige Rolle beispielsweise bei der Aufnahme von Oberflä-

chenabfluss (Abflussregulation) und Bodenpartikeln (Erosionsprävention), bei der Festlegung oder dem Transport von Nährstoffen und Pestiziden sowie im Hinblick auf die Biodiversität spielen (DOLLINGER et al. 2015, HERZON & HELENIUS 2008). Nicht zuletzt können Entwässerungsgräben als Quellen verstärkter Treibhausgasemissionen fungieren (KOSCHORRECK et al. 2020, PEACOCK et al. 2021, SCHRIER-UIJL et al. 2011) und sind daher möglicherweise auch in diesem Bereich relevant.

7.3. Nutzungshinweise und FAQ

Die Ergebnisse der hier vorgestellten Grabenerkennung sind als online-Ressource auf dem *NIBIS®-Kartenserver* veröffentlicht. Bei dem zugehörigen Datensatz handelt es sich um Rasterdaten, welche im sogenannten Tagged Image File Format (TIF) vorliegen und auf das Koordinatensystem „ETRS89 / UTM zone 32N (zE-N)“ (EPSG:4647) referenziert sind. Die Rasterzellengröße beträgt analog zum verwendeten Digitalen Geländemodell 1 x 1 m.

Bei der Betrachtung der Ergebnisse ist es ggf. sinnvoll, zunächst einen kleinen Kartenausschnitt auszuwählen, um einen besseren Eindruck der rasterbasierten Darstellung zu bekommen. Hierbei können auch kleinräumige Besonderheiten, wie beispielsweise die Unterbrechung von Entwässerungsgräben durch Auffahrten auf landwirtschaftlich genutzte Flächen, sichtbar werden. Bei der Betrachtung eines größeren Kartenausschnittes ist zu berücksichtigen, dass die modellierten Strukturen durch die Darstellung im Webbrowser möglicherweise breiter erscheinen als sie tatsächlich sind.

Nachfolgend ist ein kurzes FAQ für die Betrachtung der online-Ressource zusammengestellt.

- Die Gräben werden teilweise nicht angezeigt, woran liegt das?

Die Gräben weisen in der Regel nur eine Breite von wenigen Rasterzellen auf. Bei sehr großen Kartenausschnitten können die einzelnen Strukturen daher technisch bedingt nicht immer detailgetreu dargestellt werden. Möglicherweise muss in diesem Fall ein kleinerer Kartenausschnitt für die Betrachtung gewählt werden.

- Die Karte wirkt unruhig, es werden sehr viele Gräben angezeigt.

Je nach Region bzw. Kartenausschnitt können mitunter sehr viele Gräben auftreten. Zudem verlaufen diese nur in bestimmten Bereichen in einem gleichmäßigen „Netz“. Vielmehr orientieren sich die Gräben zu meist an den landwirtschaftlichen Nutzflächen und deren Grenzen sowie an bestehenden Straßen und Wegen.

- In der Karte werden Gräben ausgewiesen, die in der Realität keine Gräben sind.

Das Ergebnis ist nicht fehlerfrei. Es kann nicht ausgeschlossen werden, dass zum Teil auch andere Strukturen wie beispielsweise Hohlwege, natürliche Geländeeinschnitte und Gewässer sowie ins Gelände eingeschnittene schmale Straßen oder Einfahrten als Gräben erkannt werden. Des Weiteren kann es sein, dass ehemals zum Zeitpunkt der Erstellung des Digitalen Geländemodells vorhandene Gräben zwischenzeitlich überbaut, zurückgebaut oder durch Verkrautung oder Materialeintrag nicht mehr als solche zu erkennen sind. Es kann daher sinnvoll sein, für eine genaue Einschätzung entsprechende Luftbilder oder das DGM1 hinzuzuziehen.

- In der Realität gibt es Gräben, die nicht in der Karte ausgewiesen werden.

Das Ergebnis ist nicht fehlerfrei. Es kann nicht ausgeschlossen werden, dass zum Teil vorhandene Gräben nicht vom Modell erkannt werden. Dies betrifft insbesondere sehr flache Strukturen, beispielsweise an Straßen oder Wegen, und ist in gewissem Sinne beabsichtigt bzw. wurde in Kauf genommen, da ansonsten viele sonstige flache Vertiefungen im Gelände das Ergebnis verzerrt hätten. Zudem können sehr flache Gräben zwar für das Abführen von Oberflächenabfluss relevant sein, weisen jedoch allgemein keine nennenswerte Drainagewirkung auf. Des Weiteren ist es möglich, dass Gräben erst nach dem Zeitpunkt der Erstellung des Digitalen Geländemodells angelegt oder vertieft wurden. Es kann daher sinnvoll sein, für eine genaue Einschätzung entsprechende Luftbilder oder das DGM1 hinzuzuziehen.

- Ich habe Anmerkungen, Fragen oder Vorschläge für zukünftige Verbesserungen.

Für Rückfragen etc. steht die E-Mail-Adresse grundwasser@lbeg.niedersachsen.de zur Verfügung. Entsprechende Hinweise können dabei helfen, die Modellierung in Zukunft zu verbessern.

8. Dank

Die Autorin und die Autoren bedanken sich bei Martin Hoetmer sowie bei Gabriele Ertl, Eva González, Grit Griffel, Ines Hoopmann, Kai Holzgräfe, Christian Röder, Dr. Claudia Schimschal, Toni Widmer und Melanie Witthöft für die Rückmeldungen und zum Teil intensiven Plausibilitätsprüfungen im Gelände. Für technische Hinweise und fachliche Tipps bedanken wir uns bei Dr. Jonas Bostelmann, Christian Hönniger und Malte Manne sowie bei Dr. Nico Herrmann, Dr. André Kirchner und Dr. Robin Stadtmann.

9. Quellen

- BAILLY, J. S., LAGACHERIE, P., MILLIER, C., PUECH, C. & KOSUTH, P. (2008): Agrarian landscapes linear features detection from LiDAR: Application to artificial drainage networks. – *International Journal of Remote Sensing* **29** (12): 3489–3508; <https://doi.org/10.1080/01431160701469057>.
- BALADO, MARTÍNEZ-SÁNCHEZ, ARIAS & NOVO (2019): Road Environment Semantic Segmentation with Deep Learning from MLS Point Cloud Data. – *Sensors* **19** (16): 3466; <https://doi.org/10.3390/s19163466>.
- BEN-DAVID, A. (2008): Comparison of classification accuracy using Cohen's Weighted Kappa. – *Expert Systems with Applications* **34** (2): 825–832; <https://doi.org/10.1016/j.eswa.2006.10.022>.
- BREIMAN, L. (1996): Bagging predictors. – *Machine Learning* **24** (2): 123–140; <https://doi.org/10.1007/BF00058655>.
- BREIMAN, L. (2001): Random Forests. – *Machine Learning* **45** (1): 5–32; <https://doi.org/10.1023/A:1010933404324>.
- BRENNING, A., BANGS, D. & BECKER, M. (2022): RSAGA: SAGA Geoprocessing and Terrain Analysis (Version 1.4.0) [Software]. – <https://CRAN.R-project.org/package=RSAGA>.
- BROERSEN, T., PETERS, R. & LEDOUX, H. (2017): Automatic identification of watercourses in flat and engineered landscapes by computing the skeleton of a LiDAR point cloud. – *Computers & Geosciences* **106**: 171–180; <https://doi.org/10.1016/j.cageo.2017.06.003>.
- BUCHANAN, B. P., FALBO, K., SCHNEIDER, R. L., EASTON, Z. M. & WALTER, M. T. (2013): Hydrological impact of roadside ditches in an agricultural watershed in Central New York: Implications for non-point source pollutant transport. – *Hydrological Processes* **27** (17): 2422–2437; <https://doi.org/10.1002/hyp.9305>.
- CAZORZI, F., FONTANA, G. D., LUCA, A. D., SOFIA, G. & TAROLLI, P. (2013): Drainage network detection and assessment of network storage capacity in agrarian landscape. – *Hydrological Processes* **27** (4): 541–553; <https://doi.org/10.1002/hyp.9224>.
- CHEN, T. & GUESTRIN, C. (2016): XGBoost: A Scalable Tree Boosting System. – *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*: 785–794; <https://doi.org/10.1145/2939672.2939785>.
- CHEN, T., HE, T., BENESTY, M., KHOTILOVICH, V., TANG, Y., CHO, H., CHEN, K., MITCHELL, R., CANO, I., ZHOU, T., LI, M., XIE, J., LIN, M., GENG, Y., LI, Y. & YUAN, J. (2024): xgboost: Extreme Gradient Boosting (Version 1.7.8.1) [Software]; <https://CRAN.R-project.org/package=xgboost>.
- COHEN, J. (1960): A Coefficient of Agreement for Nominal Scales. – *Educational and Psychological Measurement* **20** (1): 37–46; <https://doi.org/10.1177/001316446002000104>.
- CONRAD, O., BECHTEL, B., BOCK, M., DIETRICH, H., FISCHER, E., GERLITZ, L., WEHBERG, J., WICHMANN, V. & BÖHNER, J. (2015): System for Automated Geoscientific Analyses (SAGA) v. 2.1.4. – *Geoscientific Model Development* **8** (7): 1991–2007; <https://doi.org/10.5194/gmd-8-1991-2015>.
- COVERT, I. & LEE, S.-I. (2021): Improving Kernel-SHAP: Practical Shapley Value Estimation via Linear Regression. – *arXiv:2012.01536*; <https://doi.org/10.48550/arXiv.2012.01536>.
- DOLLINGER, J., DAGÈS, C., BAILLY, J.-S., LAGACHERIE, P. & VOLTZ, M. (2015): Managing ditches for agroecological engineering of landscape. A review. – *Agronomy for Sustainable Development* **35** (3): 999–1020; <https://doi.org/10.1007/s13593-015-0301-6>.
- DU, L., MCCARTY, G. W., LI, X., ZHANG, X., RABENHORST, M. C., LANG, M. W., ZOU, Z., ZHANG, X. & HINSON, A. L. (2024): Drainage ditch network extraction from lidar data using deep convolutional neural networks in a low relief landscape. – *Journal of Hydrology* **628**: 130591; <https://doi.org/10.1016/j.jhydrol.2023.130591>.
- ERTL, G., BUG, J., ELBRACHT, J., ENGEL, N. & HERRMANN, F. (2019): Grundwasserneubildung von Niedersachsen und Bremen. Berechnungen mit dem Wasserhaushaltsmodell mGROWA18. – *GeoBerichte* **36**: 54 S., 20 Abb., 9 Tab.; Hannover (LBEG); DOI 10.48476/geoerber_36_2019.

- ERTL, G., BUG, J., GHALICHEHBAF, S., HAJATI, M.-C., HERRMANN, F., WALDECK, A., & ELBRACHT, J. (2024): Die Grundwasserneubildung und weitere Wasserhaushaltskomponenten von Niedersachsen - Berechnungen mit dem Wasserhaushaltsmodell mGROWA22. – *GeoBerichte* **51**: 66 S., 30 Abb., 3 Tab., Anh.; Hannover (LBEG); DOI 10.48476/geober_51_2024.
- FLYCKT, J., ANDERSSON, F., LAVESSON, N., NILSSON, L. & Å GREN, A. M. (2022): Detecting ditches using supervised learning on high-resolution digital elevation models. – *Expert Systems with Applications* **201**: 116961; <https://doi.org/10.1016/j.eswa.2022.116961>.
- FRIEDMAN, J. H. (1999): Stochastic Gradient Boosting. – 10 S., [Technical Report]; Stanford University.
- FRIEDMAN, J. H. (2001): Greedy function approximation: A gradient boosting machine. – *The Annals of Statistics* **29** (5): 1189–1232; <https://doi.org/10.1214/aos/1013203451>.
- GEHRT, E., BENNE, I., EVERTSBUSCH, S., KRÜGER, K. & LANGNER, S. (2021): Erläuterung zur BK 50 von Niedersachsen. – *GeoBerichte* **40**: 282 S., 125 Abb., 100 Tab.; Hannover (LBEG); DOI 10.48476/geober_40_2021.
- GEHRT, E., BUG, J. & WALDECK, A. (2019): Potenzielle Drängebiete in Niedersachsen auf Grundlage der Bodenkarte von Niedersachsen im Maßstab 1 : 50.000 (BK 50). – *Geofakten* **34**: 12 S., 7 Abb., 1 Tab.; Hannover (LBEG); DOI 10.48476/geofakt_34_1_2019.
- GRAF, M., HÖPER, H. & HAUCK-BRAMSIEPE, K. (Redaktion) (2022): Handlungsempfehlungen zur Renaturierung von Hochmooren in Niedersachsen. – *GeoBerichte* **45**: 117 S., 45 Abb., 8 Tab.; Hannover (LBEG); DOI 10.48476/geober_45_2022.
- GRAMLICH, A., STOLL, S., ALDRICH, A., STAMM, C., WALTER, T. & PRASHUHN, V. (2018): Einflüsse landwirtschaftlicher Drainage auf den Wasserhaushalt, auf Nährstoffflüsse und Schadstoffaustrag. Eine Literaturstudie. – *Agroscop Science* **73**: 53; Zürich (Agroscope).
- GRAVES, J., MOHAPATRA, R. & FLATGARD, N. (2020): Drainage Ditch Berm Delineation Using Lidar Data: A Case Study of Waseca County, Minnesota. – *Sustainability* **12** (22): 9600; <https://doi.org/10.3390/su12229600>.
- GREENWELL, B. M. & BOEHMKE, B. C. (2020): Variable Importance Plots—An Introduction to the vip Package. – *The R Journal* **12** (1): 343–366.
- GUDIM, S. (2019): Impoundment Size Index (ISI): The Use Of Potential Impoundment Size For The Characterization Of Topographic Incision In DEMs. – Masters Thesis: 73 S., University of Guelph, Ontario, Canada; <https://atrium.lib.uoguelph.ca/items/d154a52a-5c1a-464a-8e6a-fbe206059038>.
- HASTIE, T., TIBSHIRANI, S. & FRIEDMAN, H. (2009): The Elements of Statistical Learning: Data Mining, Inference, and Prediction. – Second Edition; Springer.
- HERZON, I. & HELENIUS, J. (2008): Agricultural drainage ditches, their biological importance and functioning. – *Biological Conservation* **141** (5): 1171–1183; <https://doi.org/10.1016/j.biocon.2008.03.005>.
- HIJMANS, R. J. (2023): terra: Spatial Data Analysis (Version 1.7-39) [Software]. – <https://CRAN.R-project.org/package=terra>.
- HÖNNIGER, C. (2024, Oktober): Befliegungsstrategie Niedersachsen. – Online-Präsentation.
- KOKALJ, Ž. & HESSE, R. (2017): Airborne laser scanning raster data visualization. A Guide to Good Practice. – *Prostor, kraj, čas* **14**, Online ISSN 2335-4208, ZRC SAZU, Založba ZRC; <https://doi.org/10.3986/9789612549848>.
- KOKALJ, Ž., ZAKŠEK, K. & OŠTIR, K. (2011): Application of sky-view factor for the visualisation of historic landscape features in lidar-derived relief models. – *Antiquity* **85** (327): 263–273; <https://doi.org/10.1017/S0003598X00067594>.
- KOSCHORRECK, M., DOWNING, A. S., HEJZLAR, J., MARCÉ, R., LAAS, A., ARNDT, W. G., KELLER, P. S., SMOLDERS, A. J. P., VAN DIJK, G. & KOSTEN, S. (2020): Hidden treasures: Human-made aquatic ecosystems harbour unexplored opportunities. – *Ambio* **49**: 531–540; <https://doi.org/10.1007/s13280-019-01199-6>.
- KUHN, M. & WICKHAM, H. (2023): Tidymodels: A collection of packages for modeling and machine learning using tidyverse principles (Version 1.1.0) [Software]. – <https://www.tidymodels.org>.

- LANDIS, J. R. & KOCH, G. G. (1977): The Measurement of Observer Agreement for Categorical Data. – *Biometrics* **33** (1): 159–174; <https://doi.org/10.2307/2529310>.
- LGLN – LANDESAMT FÜR GEOINFORMATION UND LANDESVERMESSUNG NIEDERSACHSEN (2024a): Digitale Orthophotos (ATKIS-DOP). – [© GeoBasis-DE/LGLN 2024]; <https://www.lgln.niedersachsen.de/>.
- LGLN – LANDESAMT FÜR GEOINFORMATION UND LANDESVERMESSUNG NIEDERSACHSEN (2024b): Digitales Geländemodell (DGM1). – [© GeoBasis-DE/LGLN 2024]; <https://www.lgln.niedersachsen.de/>.
- LGLN – LANDESAMT FÜR GEOINFORMATION UND LANDESVERMESSUNG NIEDERSACHSEN (2024c): Digitales Landschaftsmodell (Basis-DLM). – [© GeoBasis-DE/LGLN 2024]; <https://www.lgln.niedersachsen.de/>.
- LIDBERG, W., PAUL, S. S., WESTPHAL, F., RICHTER, K. F., LAVESSON, N., MELNIKS, R., IVANOV, J., CIESIELSKI, M., LEINONEN, A. & ÅGREN, A. M. (2023): Mapping Drainage Ditches in Forested Landscapes Using Deep Learning and Aerial Laser Scanning. – *Journal of Irrigation and Drainage Engineering* **149** (3); <https://doi.org/10.1061/JIDEDH.IRENG-9796>.
- LINDSAY, J. B. (2016): Whitebox GAT: A case study in geomorphometric analysis. – *Computers & Geosciences* **95**: 75–84; <https://doi.org/10.1016/j.cageo.2016.07.003>.
- LUNDBERG, S. M., ERION, G., CHEN, H., DEGRAVE, A., PRUTKIN, J. M., NAIR, B., KATZ, R., HIMMELFARB, J., BANSAL, N. & LEE, S.-I. (2019): Explainable AI for Trees: From Local Explanations to Global Understanding. – *arXiv:1905.04610*; <https://doi.org/10.48550/arXiv.1905.04610>.
- LUNDBERG, S. M. & LEE, S.-I. (2017): A Unified Approach to Interpreting Model Predictions. – In: I. GUYON, U. V. LUXBURG, S. BENGIO, H. WALLACH, R. FERGUS, S. VISHWANATHAN & R. GARNETT (Hrsg.): *Advances in Neural Information Processing Systems* **30**: 10 S.; Curran Associates, Inc.; https://proceedings.neurips.cc/paper_files/paper/2017/file/8a20a8621978632d76c43dfd28b67767-Paper.pdf.
- MAYER, M. (2024): shapviz: SHAP Visualizations (Version 0.9.3) [Software]. – <https://CRAN.R-project.org/package=shapviz>.
- MAYER, M. & WATSON, D. (2024): kernelshap: Kernel SHAP (Version 0.6.0) [Software]. – <https://CRAN.R-project.org/package=kernelshap>.
- MOLNAR, C. (2022): Interpretable machine learning: A guide for making black box models explainable. – (Second edition); Christoph Molnar.
- MOLNAR, C. (2023). Interpreting Machine Learning Models With SHAP - A Guide With Python Examples And Theory On Shapley Values. – Christoph Molnar.
- OLSEN, L. R. & ZACHARIAE, H. B. (2024): cvms: Cross-Validation for Model Selection (Version 1.6.1) [Software]. – <https://CRAN.R-project.org/package=cvms>.
- PASSALACQUA, P., BELMONT, P. & FOUFOULA-GEORGIOU, E. (2012): Automatic geomorphic feature extraction from lidar in flat and engineered landscapes. – *Water Resources Research* **48** (3); <https://doi.org/10.1029/2011WR010958>.
- PEACOCK, M., AUDET, J., BASTVIKEN, D., COOK, S., EVANS, C. D., GRINHAM, A., HOLGERSON, M. A., HÖGBOM, L., PICKARD, A. E., ZIELINSKI, P. & FUTTER, M. N. (2021): Small artificial waterbodies are widespread and persistent emitters of methane and carbon dioxide. – *Global Change Biology* **27** (20): 5109–5123; <https://doi.org/10.1111/gcb.15762>.
- POSIT TEAM (2023): RStudio: Integrated Development Environment for R (Version 2023.09.1+494) [Software]. – Posit Software, PBC; <http://www.posit.co/>.
- QIAN, T., SHEN, D., XI, C., CHEN, J. & WANG, J. (2018): Extracting Farmland Features from Lidar-Derived DEM for Improving Flood Plain Delineation. – *Water* **10** (3): 252; <https://doi.org/10.3390/w10030252>.
- R CORE TEAM (2023): R: A Language and Environment for Statistical Computing (Version 4.3.2) [Software]. – R Foundation for Statistical Computing; <https://www.R-project.org/>.
- RAPINEL, S., HUBERT-MOY, L., CLÉMENT, B., NABUCET, J. & CUDENNEC, C. (2015): Ditch network extraction and hydrogeomorphological characterization using LiDAR-derived DTM in wetlands. – *Hydrology Research* **46** (2): 276–290; <https://doi.org/10.2166/nh.2013.121>.

- ROBB, C., PICKARD, A., WILLIAMSON, J. L., FITCH, A. & EVANS, C. (2023): Peat Drainage Ditch Mapping from Aerial Imagery Using a Convolutional Neural Network. – *Remote Sensing* **15** (2): 499; <https://doi.org/10.3390/rs15020499>.
- RÖDER, C. (2022): Feststofftransporte über Gewässer im Einzugsgebiet des Dümmer Sees. – [Präsentation], 14. BVB-Jahrestagung am 20.09.2022; Hannover.
- RÖDER, C. & SCHÄFER, W. (2015): Diffuse Phosphoreinträge im Einzugsgebiet des Dümmer - Quellen und Transportpfade in die Vorfluter, Ableitung von Maßnahmen. – 37 S.; Hannover (LBEG).
- ROELEN, J., HÖFLE, B., DONDEYNE, S., VAN ORSHOVEN, J. & DIELS, J. (2018): Drainage ditch extraction from airborne LiDAR point clouds. – *ISPRS Journal of Photogrammetry and Remote Sensing* **146**: 409–420; <https://doi.org/10.1016/j.isprsjprs.2018.10.014>.
- ROELEN, J., ROSIER, I., DONDEYNE, S., VAN ORSHOVEN, J. & DIELS, J. (2018): Extracting drainage networks and their connectivity using LIDAR data. – *Hydrological Processes* **32** (8): 1026–1037; <https://doi.org/10.1002/hyp.11472>.
- SCHRIER-UIJL, A. P., VERAART, A. J., LEFFELAAR, P. A., BERENDSE, F. & VEENENDAAL, E. M. (2011): Release of CO₂ and CH₄ from lakes and drainage ditches in temperate wetlands. – *Biogeochemistry* **102**: 265–279; <https://doi.org/10.1007/s10533-010-9440-7>.
- SHAPLEY, L. S. (1953): A Value for n-Person Games. – In: *Contributions to the Theory of Games II*: 307–317; Princeton University Press.
- SLA – SERVICEZENTRUM LANDENTWICKLUNG UND AGRARFÖRDERUNG (2022): Feldblöcke für Niedersachsen. – [© ML/SLA Niedersachsen/CC BY 4.0]; <https://www.sla.niedersachsen.de>.
- STANISLAWSKI, L., BROCKMEYER, T. & SHAVERS, E. (2018): Automated road breaching to enhance extraction of natural drainage networks from elevation models through deep learning. – *The International Archives of the Photogrammetry, Remote Sensing and Spatial Information Sciences* **XLII-4**: 597–601; <https://doi.org/10.5194/isprs-archives-XLII-4-597-2018>.
- TETZLAFF, B., KUHR, P. & WENDLAND, F. (2008): Ein neues Verfahren zur differenzierten Ableitung von Dränflächenkarten für den mittleren Maßstabsbereich auf Basis von Luftbildern und Geodaten. – *Hydrologie und Wasserbewirtschaftung* **52** (1): 9–18.
- TOMASI, C. & MANDUCHI, R. (1998): Bilateral filtering for gray and color images. – *Sixth International Conference on Computer Vision (IEEE Cat. No.98CH36271)*: 839–846; <https://users.soe.ucsc.edu/~manduchi/Papers/ICCV98.pdf>.
- WICKHAM, H., AVERICK, M., BRYAN, J., CHANG, W., MCGOWAN, L. D., FRANÇOIS, R., GROLEMUND, G., HAYES, A., HENRY, L., HESTER, J., KUHN, M., PEDERSEN, T. L., MILLER, E., BACHE, S. M., MÜLLER, K., OOMS, J., ROBINSON, D., SEIDEL, D. P., SPINU, V., TAKAHASHI, K., VAUGHAN, D., WILKE, C., WOO, K. & YUTANI, H. (2019): Welcome to the tidyverse. – *Journal of Open Source Software* **4** (43): 1686; <https://doi.org/10.21105/joss.01686>.
- WRIGHT, M. N. & ZIEGLER, A. (2017): ranger: A Fast Implementation of Random Forests for High Dimensional Data in C++ and R. – *Journal of Statistical Software* **77** (1): 1–17; <https://doi.org/10.18637/jss.v077.i01>.
- WU, Q. & BROWN, A. (2022): „whitebox“: „WhiteboxTools“ R Frontend (Version 2.2.0) [Software]. – <https://CRAN.R-project.org/package=whitebox>.
- YOKOYAMA, R., SHIRASAWA, M. & PIKE, R. J. (2002): Visualizing Topography by Openness: A New Application of Image Processing to Digital Elevation Models. – *Photogrammetric Engineering & Remote Sensing* **68** (3): 251–266.

Anhang

A.1 Interpretation des Kappa-Indices

Für die Interpretation des Kappa-Indices wurde auf das Bewertungsschema nach LANDIS & KOCH (1977) zurückgegriffen. Für den vorliegenden Text wurden folgende Übersetzungen gewählt:

Kappa-Index	Strength of Agreement	Stärke der Übereinstimmung
< 0,00	poor	keine
0,00 – 0,20	slight	sehr gering
0,21 – 0,40	fair	gering
0,41 – 0,60	moderate	mittel
0,61 – 0,80	substantial	hoch
0,81 – 1,00	almost perfect	sehr hoch

Autorenschaft

- Jost Wessels

Landesamt für Bergbau, Energie und Geologie,
Referat L 2.5 Hydrogeologische Grundlagen,
Stilleweg 2,
30655 Hannover.

- Dr. Mithra-Christin Hajati

<https://orcid.org/0000-0002-9732-0021>

Landesamt für Bergbau, Energie und Geologie,
Referat L 2.5 Hydrogeologische Grundlagen,
Stilleweg 2,
30655 Hannover.

- Dr. Jörg Elbracht

<https://orcid.org/0000-0002-2045-9180>

Landesamt für Bergbau, Energie und Geologie,
Referat L 2.5 Hydrogeologische Grundlagen,
Stilleweg 2,
30655 Hannover.

ISSN 1864 – 7529