

Zum variationslinguistischen Verhältnis von Stadt und Land. Ein Beitrag zu einer sprachlichen Raumtopologie am Beispiel Österreichs

Pickl, Simon; Pröll, Simon; Elspaß, Stephan

Zum variationslinguistischen Verhältnis von Stadt und Land. Ein Beitrag zu einer sprachlichen Raumtopologie am Beispiel Österreichs

Linguistik online, vol. 110, núm. 5, 2021

Universität Bern, Suiza

Disponible en: <https://www.redalyc.org/articulo.oa?id=664572262013>

DOI: <https://doi.org/10.13092/lo.110.8147>



Esta obra está bajo una Licencia Creative Commons Atribución 3.0 Internacional.

Zum variationslinguistischen Verhältnis von Stadt und Land. Ein Beitrag zu einer sprachlichen Raumtopologie am Beispiel Österreichs

Simon Pickl

Universität Salzburg, Australia

Simon Pröll

Ludwig-Maximilians-Universität München, Alemania

Stephan Elspaß

Universität Salzburg, Austria

Linguistik online, vol. 110, núm. 5, 2021

Universität Bern, Suiza

DOI: <https://doi.org/10.13092/lo.110.8147>

Redalyc: <https://www.redalyc.org/articulo.oa?id=664572262013>

Abstract: In this paper, we explore the geolinguistic relationship between urban and rural areas through the conceptualisation and modeling of spatial topologies. Geolinguistic topologies concern the structure of the mutual linguistic relationship between localities. They can be defined either deductively or on the basis of empirical data and represent the linguistic similarities or distances between localities. We operationalise and apply several such topological models to Austrian data from the Atlas zur deutschen Alltagssprache (AdA), a linguistic atlas documenting colloquial German using crowd-sourcing methods. The results are evaluated on the basis of statistical examination and of visualisations of the topological relationships predicted by the models. It is confirmed that linguistic similarity is determined both by geographical distance and by the distribution of population, but the exact relationship is complex: Not only do smaller geographic distances on the one hand and higher population numbers on the other hand bring about increased linguistic similarity; the relevance of these two factors for linguistic similarity varies with population size, too, such that linguistic relationships between cities are determined more by their size and less by their distance, while for smaller locations the opposite is true. Hence, no single topological model can be identified as superior; instead, the individual models emphasise different aspects of the linguistic relationship between urban and regional language usage.

1 Einleitung: Raumtopologien

Es ist seit Langem bekannt, dass Städte im Sprachraum¹ eine besondere Stellung einnehmen (cf. u. a. Becker 1942; Bach 1950; Frings 1956): Sie verhalten sich anders als kleinere Orte, für deren sprachgeographische Stellung in erster Linie ihre räumliche Position im Verhältnis zu ihren unmittelbaren Nachbarn maßgeblich ist. Städte sind anders von Diffusionsprozessen betroffen und bedingen so räumliche Strukturen wie Trichter oder Enklaven: Neuerungen springen häufig zunächst von Stadt zu Stadt, bevor sie sich von dort aus auf deren Umland ausbreiten. Städte sind sich sprachlich oft ähnlicher, als nur aufgrund ihrer relativen geographischen Lage zu erwarten wäre; sie zeigen „insgesamt keine geographisch kohärente Verteilung“ (Pickl 2013b: 23), sondern treten „als eigene, räumlich diskontinuierliche Struktur hervor“ (Pröll 2015: 160)

und zeigen insgesamt eine stärkere Tendenz zu standardsprachlichen Repertoires als ländliche Gebiete (cf. Hernández-Campoy 2003; Pickl 2013a; Pickl et al. 2019).

Die Sonderstellung des Sprachgebrauchs in Städten² geht teils auf die sprachliche Vertikalisierung des Varietätenspektrums der frühen Neuzeit (Reichmann 1988) zurück (cf. Pickl/Elspaß 2019), teils auf Migration und Sprachkontakt (cf. Kerswill 1993; Britain 2012a; Newerla 2013), teils sind sie der sozialen Fragmentierung und der damit verbundenen variationslinguistischen Differenzierung geschuldet (Pröll/Elspaß/Pickl im Druck). Abgesehen von solchen innerstädtischen Entwicklungen bestehen aber auch Besonderheiten im sprachlichen Verhältnis von Städten zu ihrer näheren und weiteren Umgebung – sowohl zu eher dörflich geprägten Orten als auch zu anderen Städten unterschiedlicher Größe –, die sich aus verschiedenen Blickwinkeln betrachten lassen: Diese Besonderheiten bewirken u. a., dass das klassische, räumlich homogene Wellenmodell sprachlichen Wandels nur sehr eingeschränkt auf Städte anwendbar ist (siehe Abschnitt 4.1). Einerseits ist etwa der sprachliche Einfluss von Städten auf ihr ländliches Umfeld größer als umgekehrt, i. e. es handelt sich nicht um ein symmetrisches Verhältnis zwischen Städten und kleineren Orten; zum anderen ist die Tendenz zur sprachlichen Annäherung zwischen Städten oft größer als zwischen kleineren Orten bei gleicher Distanz. Beide Größen – Einfluss und sprachliche Ähnlichkeit – sind eng miteinander verbunden (siehe Abschnitt 2.3.1), denn größerer Einfluss bedingt größere Ähnlichkeit (cf. u. a. Nerbonne/Heeringa 2007). Entsprechende Verhältnisse können in Form eines topologischen Modells konzeptualisiert werden, in dem die Lagebeziehungen unter Orten durch ihre gegenseitige sprachliche Ähnlichkeit oder ihren Einfluss aufeinander definiert sind. Grundlage für solche Modelle ist die sprachliche Distanz zwischen Orten,³ die anders als ihre geographische Distanz einerseits als Kehroder Komplementärwert ihrer sprachlichen Ähnlichkeit definiert werden kann (cf. u. a. Goebel 1984) und andererseits als Faktor für die Schätzung ihrer gegenseitigen Einflusswahrscheinlichkeit verwendet wird (cf. u. a. Trudgill 1974: 234–235). Anders als beim rein geographischen, euklidischen Raum, der als Koordinatensystem einen bloßen kartesischen „Container“ (cf. u. a. Wardenga 2002; de Langer/Nipper 2018) für räumliche Objekte darstellt, sind sprachliche Distanzen nicht euklidisch, sondern definieren als paarweise Relationen selbst sprachliche Räume „als Systeme von Lagebeziehungen materieller Objekte“ (Wardenga 2002: 47), i. e. als Topologien. Löw (2001: 131) versteht „Raum als eine relationale (An)Ordnung von Körpern, welche unaufhörlich in Bewegung sind, wodurch sich die (An)Ordnung selbst ständig verändert“. In solchen sprachlichen Raumtopologien manifestiert sich die besondere Stellung der Städte gegenüber ländlich geprägten Orten.

Die Frage, wie solche sprachräumlichen Topologien strukturiert sind, ist bis jetzt unseres Wissens nicht gestellt, geschweige denn explizit diskutiert worden. Es gibt jedoch verschiedene Ansätze, sprachliche Relationen zwischen verschiedenen Orten auf der Grundlage

verschiedener Faktoren zu quantifizieren, etwa im Rahmen von Diffusionsmodellen oder von dialektometrischen Analysemethoden. Im ersten, deduktiv orientierten Fall werden die Beziehungen zwischen Orten auf der Grundlage externer Variablen – im Normalfall geographische Distanz, Bevölkerungszahlen sowie manchmal Reisestrecke bzw. -dauer, insofern diese von der Luftlinie abweichen (cf. Pickl 2013a: 61; Pröll 2015: 50) – modelliert; im zweiten, induktiv geprägten Fall werden sprachliche Distanzen aufgrund empirischer Sprachdaten, wie sie etwa in Sprachatlanten dokumentiert sind, durch Aggregation errechnet. In diesem Beitrag sollen beide Herangehensweisen zusammengeführt werden, indem eine kleine Zahl solcher *a priori* formulierten Modelle (die in Abschnitt 2 kurz charakterisiert werden) anhand von empirischen Daten evaluiert werden, indem ihre Erklärungs- bzw. Vorhersagequalität bei der Anwendung auf reale Daten überprüft wird (Abschnitte 4.1–4.3). Daran anknüpfend wird ein empirisch begründetes, deduktiv ermitteltes Alternativmodell vorgestellt, dessen Funktionsweise keine explizite Theorie zugrunde liegt (Abschnitt 4.4). In einem letzten Schritt werden die verschiedenen Modelle auf ihre Abbildungsadäquatheit sowie ihre theoretische Fundierung hin verglichen (Abschnitt 5), bevor die Untersuchung mit einem kurzen Fazit schließt (Abschnitt 6.) Die Untersuchung versteht sich als Fallstudie, die die Angemessenheit verschiedener Modelle in einem konkreten Anwendungsfall beurteilt und dabei keine Allgemeingültigkeit der Ergebnisse beansprucht, aber dennoch prinzipiell verallgemeinerbare Aussagen beabsichtigt. Dabei ist insbesondere von Interesse, wie die einzelnen Modelle die sprachliche Stellung von Städten in sprachlichen Raumtopologien in Relation zur regionalen Sprachvariation erfassen, und welche Schlüsse daraus für das sprachliche Verhältnis von urbanen und ruralen Räumen zu ziehen sind.

Das Anwendungsbeispiel in diesem Beitrag betrifft die Alltagssprache in Österreich, wie sie im *Atlas zur deutschen Alltagssprache* (AdA, Elspaß/Möller 2003–) dokumentiert ist (siehe im Detail Abschnitt 3). Datenbasis stellen die AdA-Erhebungsrunden 8–10 dar (siehe Abschnitt 3; cf. auch Pickl et al. 2019: 51).

2 Topologische Modelle des Sprachraums

Der Nutzen der Modellierung von sprachräumlichen Topologien besteht darin, die relativen Lagebeziehungen von Siedlungspunkten so zu erfassen, dass dadurch diejenigen Beziehungen (bzw. sprachlichen Ähnlichkeiten) zwischen diesen Punkten möglichst gut abgebildet bzw. vorhergesagt werden, die für Diffusionsbewegungen relevant sind.⁴ Im einfachsten Fall – dem Wellenmodell – sind diese zugrundeliegenden Relationen identisch mit den rein geographischen Distanzen. Dieser „Urtyp“ der Diffusionsmodelle operiert auf der Grundlage des starren und homogenen geographischen Raumes, in dem Städte und die Bevölkerungsverteilung keine eigene Rolle spielen. Folglich haben in diesem Modell näher beieinander liegende Orte ungeachtet ihrer

Größe eine große Wahrscheinlichkeit, dieselben sprachlichen Formen aufzuweisen, während weiter entfernt voneinander liegende Orte eher dazu tendieren, sprachliche Unähnlichkeiten zu entwickeln – ein Strukturprinzip, das im Grundsatz oft empirisch nachgewiesen wurde, in der Dialektologie aber überraschend spät explizit formuliert wurde und heute als „Fundamental Dialectological Postulate“ bekannt ist (Nerbonne/Kleiweg 2007, 154): „Geographically proximate varieties tend to be more similar than distant ones“.

Bereits einige Jahrzehnte zuvor hat Séguy (1971) einen konkreteren, empirisch begründeten Beitrag zur Relation zwischen sprachlichen und räumlichen Distanzen vorgelegt (jedoch nicht im Rahmen eines Diffusionsprozesses): Er hat erstmals sprachliche Unähnlichkeit in Abhängigkeit der räumlichen Entfernung – auf der Basis mehrerer Sprachatlanten – dargestellt, und dafür einen logarithmischen Zusammenhang identifiziert. (Zunächst steigt die Unähnlichkeit stark dann, dann stellt sich Sättigung ein, da bereits bestehenden Unterschiede mit wachsender Distanz nur durch andere Unterschiede ersetzt werden, die dieselbe Unähnlichkeit aufweisen.) Ein solcher Zusammenhang wurde seitdem in vielen Studien bestätigt (cf. u. a. Heeringa/Nerbonne 2001; Nerbonne/Heeringa 2007: 288f.; Nerbonne 2010; Pickl 2013a: 101–102; Pickl et al. 2014) und ist heute als „Séguy’s curve“ bekannt (cf. Nerbonne 2010; Pickl/Pröll 2019: 862) (siehe Abbildung 1).

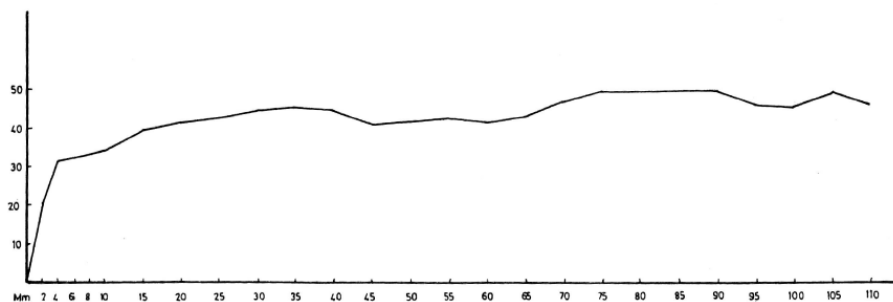


Abbildung 1

Die Séguy-Kurve (Séguy 1971: 349). Die x-Achse repräsentiert die geographische Distanz, die y-Achse die linguistische Unähnlichkeit

Der von Séguy beobachtete Zusammenhang betrifft die linguistischen Distanzen allein in Abhängigkeit von geographischen Distanzen, dem geographischen Raum, der in den Worten Hards (2008: 285) die „Mager-, Schwund-, ja Nullstufe“ des (Sprach-)Raums darstellt (cf. Pickl 2013a: 58). Daneben wurde auch versucht, sprachliche Räume topologisch, i. e. durch anderweitige paarweise Distanzbeziehungen zwischen Orten, darzustellen, wodurch der geographische Raum überformt und unter Hinzuziehung weiterer relevanter Parameter bzw. Faktoren zu einem möglichst realistischen Modell sprachlicher Räume ausgebaut wird. Als möglicher „Ersatz“ für den geographischen Abstand wurden häufig (v. a. historische) Reisedistanzen in Erwägung gezogen – mit unterschiedlichen Ergebnissen.⁵

Es sind viele weitere potentielle Faktoren denkbar, die man sinnvoll in eine Modellierung eines linguistischen Distanzraums einbeziehen könnte, darunter „[p]hysical, social, and perceptual factors (mountains, marshes, motorways, lack of roads or public transport, employment blackspots, shopping malls, xenophobia, or external negative perceptions of place“ (Britain 2002: 609) – kurz: wie auch immer geartete potentielle Barrieren und/oder Verstärker von sprachlichem Austausch (Pickl 2013a: 141–157; 2017); diese sind jedoch ungleich schwieriger systematisch zu berücksichtigen und verlangen starke *a priori*-Annahmen in der (deduktiven) Modellbildung, die im Einzelfall nicht einfach zu begründen sind. Als wichtigste Größe neben der geographischen Distanz hat sich die Bevölkerungsverteilung erwiesen, die das Gefälle zwischen urbanen und ruralen Räumen abbildet und in verschiedener Form in die Modellierung von linguistischen Abständen einfließen kann.

Neben dem bereits oben erwähnten Wellenmodell werden komplexere Diffusionsmodelle unterschieden, die in irgendeiner Form die Größe von Orten berücksichtigen:

There are three central hypotheses about how this progression takes place. The wave model (sometimes referred to as the contagion model) predicts that change spreads regularly outward from a central point (e. g. Bloomfield 1933, Trudgill 1986, Bailey et al. 1993, Labov 2003, Britain 2004, 2010). The gravity model predicts that change spreads based on population size and that progression need not be regular across geographic space (e. g. Trudgill 1974, 1986). In hierarchical models, like the cascade model, change moves from the largest to the next largest city in a predictable order within regions (e. g. Boberg 2000).
(Tagliamonte/D’Arcy/Rodríguez Louro 2016: 826)

Im Folgenden werden die einschlägigen Modelle kurz anhand konstruierter Beispiele charakterisiert.

2.1 Wellenmodell

Betrachten wir zunächst noch einmal das Wellenmodell (cf. Pickl 2013a: 51–54 für eine Diskussion der Annahmen, Grundprinzipien und soziolinguistischen Konzeptualisierung des Wellenmodells) näher, und zwar mit Blick auf das dafür maßgebliche räumliche Modell. Hierfür wurden 24 Ortspunkte mit unterschiedlichen Bevölkerungszahlen entlang einer Linie angelegt, die den zweidimensionalen geographischen Raum repräsentiert (siehe Abbildung 2); die Einheiten dieser Achse kann man sich als Kilometer vorstellen. Die Blasengröße signalisiert hier die Bevölkerungszahl (der kleinste Ort wurde mit 270, der größte mit 10 000 Einwohnerinnen und Einwohnern angelegt; siehe auch Abbildung 3). Der Farbton gibt die Position entlang der geographischen Achse wieder und dient der ungefähren Einordnung bzw. Identifikation der Orte bei den nachfolgenden Abbildungen.

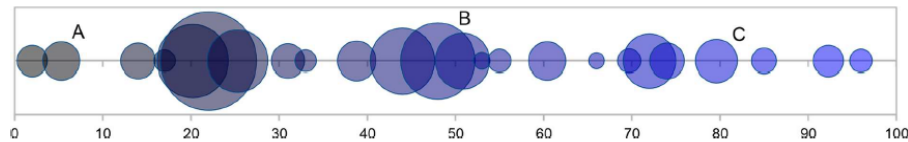


Abbildung 2
Topologische Struktur des Wellenmodells (Beispiel)

Topologisch relevant ist in dieser Darstellung ausschließlich die Position entlang der geographischen Achse. Die drei beispielhaft ausgewählten Orte A, B und C sind horizontal verbunden, i. e. die kürzeste Strecke zwischen ihnen geht über alle Orte, die geographisch zwischen ihnen liegen. Bei Diffusionen kommen dementsprechend etwa Innovationen, die sich von A ausbreiten, nur dann in C an, wenn sie vorher bereits B erfasst haben. Daher ist sprachlich mit einem homogenen Kontinuum zu rechnen: Das „Fundamental Dialectological Postulate“ gilt uneingeschränkt.

2.2 Hierarchisches Modell

Diffusionstypen, die von der rein geographisch bedingten Ausbreitung abweichen, oder Raumbilder, die nicht mit dem Fundamental Dialectological Postulate übereinstimmen, vermag das Wellenmodell nicht zu erklären, da sie zusätzlich von der Bevölkerungsgröße abhängig sind. Das hierarchische Modell (auch bekannt als „urban hierarchy model“ oder als „cascade model“) stützen sich dagegen vor allem auf die Größenrelationen der Orte:

According to this concept, innovations begin in central places, which serve as the focal points for diffusion across the landscape. Rather than spreading uniformly across the landscape, diffusion begins in large cities like London, moving then to smaller cities and so on down the hierarchy. The term “cascade diffusion” is used for changes that move downward specifically from larger to smaller cities.

(Bailey et al. 1993: 361)

Wenn wir die Orte aus Abbildung 2 zusätzlich zur horizontalen, geographischen Achse auf einer vertikalen -Achse anordnen, die für die Bevölkerungsgröße der Orte steht, ergibt sich eine Konstellation wie in Abbildung 3:

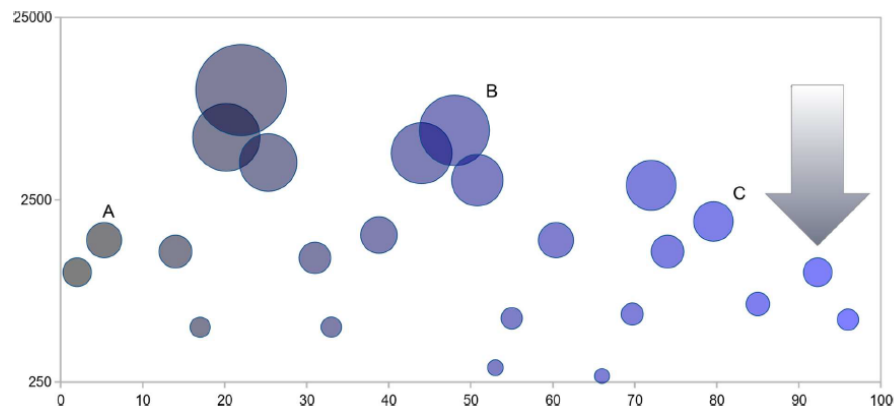


Abbildung 3
Topologische Struktur des hierarchischen Modells (Beispiel)

Diese Darstellung visualisiert zunächst folgende Vorstellung: Je bevölkerungsreicher Orte sind, umso mehr sind sie aus der Horizontalen herausgelöst. Maßgeblich für die Lagebeziehung der Orte untereinander ist im hierarchischen Modell vor allem die Vertikale, der Horizontalen kommt allenfalls eine untergeordnete Rolle zu:

[T]he urban hierarchy model [...] suggests that innovations spread down an urban hierarchy from metropolis to city to town to village to countryside. The rationale for this model is that transportation networks tend to link urban with urban, and the socioeconomic and consumer infrastructure tends to be based in and oriented towards urban centres, so that while distance plays some role, interaction between urban centres is likely to be greater, and therefore a more frequent and effective channel for innovation transmission, than between urban and rural [...].

(Britain 2016: 226)

Insofern sind die Beziehungen in einem solchen Modell nicht nur asymmetrisch, sondern klar gerichtet: Größere Orte beeinflussen kleinere, nicht umgekehrt. Die größeren Orte sind für die kleineren sehr präsent, die kleineren für die größeren jedoch praktisch bedeutungslos. Dabei ist nicht generell spezifiziert, als wie groß der Einfluss der räumlichen Distanz tatsächlich angenommen wird (i. e. wie relevant die horizontalen Relationen entlang der .-Achse in Abbildung 3 sind); in Reinform stützen sich hierarchische Modelle allein auf vertikale Relationen.

Wenn wir in Abbildung 3 wieder dieselben Orte (A, B und C) betrachten, wird klar, dass sich im Rahmen dieses Ansatzes sprachliche Neuerungen in diesen Orten nur durchsetzen können, wenn sie dort ihren Ursprung nehmen oder aus den jeweils größeren Orten kommen. Eine Diffusion unter etwa gleich großen kleineren Orten spielt in diesem Modell praktisch keine Rolle. Vor allem sieht es nicht vor, dass kleinere Orte größere beeinflussen.⁶

2.3 Interaktionale Modelle

Einen Versuch, Fälle zu erfassen, bei denen sich Neuerungen nicht homogen im Raum verbreiten, deren Diffusionsmuster sich aber auch nicht allein aufgrund einer urbanen Hierarchie erklären lässt, stellen interaktionale Modelle dar, bei denen paarweise Beziehungen zwischen Orten errechnet werden, die ihre Interaktion untereinander quantifizieren und damit ihren sprachlichen Austausch bzw. im Resultat ihre sprachliche Ähnlichkeit modellieren. Das bekannteste interaktionale Modell ist das Gravitationsmodell nach Trudgill (1974), das oft als Implementierung des hierarchischen Modells gesehen wird (cf. z. B. Bailey et al. 1993: 361; Nerbonne 2010: 3822).⁷ „The crucial difference between these models is the lack of importance of physical distance between two speech communities in the cascade model.“ (Denis 2009: 37) Aus diesem Grund, sowie wegen seiner expliziten Konzeptualisierung der sprachlichen Interaktion zwischen zwei Orten und zur besseren Kontrastierung von der stärkeren Betonung vertikaler Relationen, wird das Gravitationsmodell hier zu den interaktionalen Modellen gezählt. In dieser Hinsicht dient hier die euklidische Distanz im Gegensatz zum hierarchischen Modell als elementare Berechnungsgrundlage, aus der zusammen mit den jeweiligen Bevölkerungszahlen (und gegebenenfalls weiteren Kenngrößen) relationale Werte ermittelt werden, um so die Interaktion als Grundlage für die sprachliche Ähnlichkeit zu schätzen.

Eine Eigenschaft interaktionaler Modelle ist, dass die paarweisen Relationen, die als Maß für die sprachliche Unähnlichkeit fungieren, zwischen großen Orten durch ihre Bevölkerungs-„Massen“ gegenüber ihren geographischen Distanzen reduziert erscheinen, bevölkerungsstarke Orte werden also als „näher“ konzeptionalisiert. Dasselbe gilt – in eingeschränktem Ausmaß – für die Relationen zwischen großen und kleineren Orten, während die Relationen zwischen kleineren Orten von solchen Anziehungserscheinungen weitgehend unberührt bleiben. Die Analogie dieser „Anziehungskraft“ zwischen Orten abhängig von ihrer Größe ist die Grundlage für die metaphorische Bezeichnung von Trudgills Modell als Gravitationsmodell. Die Relationen zwischen Orten werden als Interaktionen mithilfe von Trudgills Formel aus den geographischen Distanzen und den jeweiligen Bevölkerungszahlen berechnet. Die resultierenden Werte sind die Basis für eine Topologie, i. e. sie beschreiben die relationale Anordnung der Orte zueinander als Modellierung ihrer gegenseitigen sprachlichen Beziehungen und definieren so einen topologischen Raum. Sie sind jedoch nicht euklidisch oder metrisch, deshalb lassen sie sich nicht ohne Verzerrungen in einem zweier oder dreidimensionalen Raum abbilden. Visuelle Darstellungen wie die in Abbildung 4 sind also der Anschaulichkeit geschuldete Vereinfachungen, Projektionen eines mehrdimensionalen topologischen Raums in die zweidimensionale Fläche. Abbildung 4 veranschaulicht die Topologien interaktionaler Modelle im Vergleich zu denen des Wellen- und des hierarchischen

Modells prinzipiell, i. e. ohne sich auf eine bestimmte Implementierung des interaktionalen Ansatzes zu stützen (siehe hierzu 2.3.1–2.3.2).

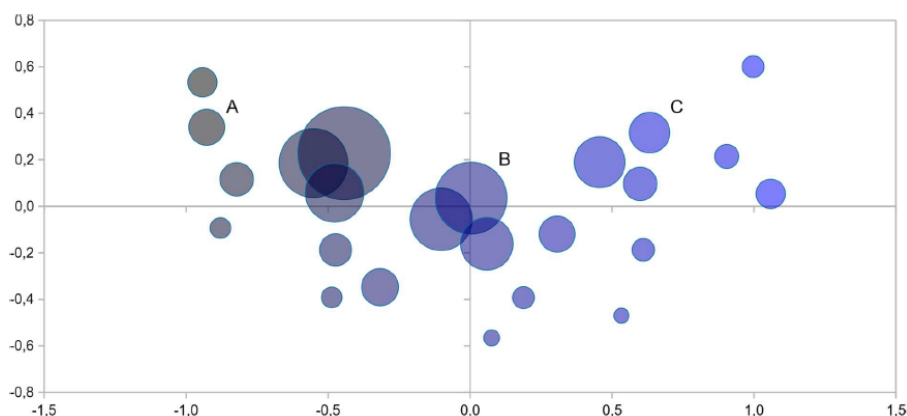


Abbildung 4
Topologische Struktur von interaktionalen Modellen (Beispiel)

Im Vergleich zum hierarchischen Modell (siehe Abbildung 3) zeigt sich hier, dass bevölkerungsreiche Orte nicht nur aus der horizontalen Achse herausgelöst sind; sie sind sich erstens auch gegenseitig näher, und zweitens haben sie dabei auch in stärkerem Maß kleinere Orte um sich versammelt. Die größeren Orte sind so näher zusammengedrückt, während kleine Orte weitgehend ihre lokalen Entfernungen untereinander zwar bewahren; sie sind jedoch an ihre jeweiligen Zentralorte gebunden und richten sich an ihnen aus; dadurch wird die bisherige horizontale Achse „gekrümmt“: Sie biegt sich um die zentrale Ansammlung an städtischen Zentren.

Mit Blick auf die oben herausgegriffenen Beispielpunkte A, B und C ergeben sich folgende Konsequenzen: Sie sind nach wie vor entlang der (nun gekrümmten) horizontalen Achse untereinander verbunden; die Verbindung über die Zentralorte ist aber kürzer als ihre horizontale Verbindung. Um von A nach C zu gelangen, stellt B also Teil einer Abkürzung dar; dennoch ist, anders als beim Kaskadenmodell, eine „horizontale“ Verbindung möglich.

2.3.1 Gravitationsmodell

Der interaktionale Ansatz einer sprachlichen Topologie ist unseres Wissens der einzige, für den es Operationalisierungen im Sinne ausformulierter, anwendbarer Modelle gibt. Die genaue Form der topologischen Distanzen und damit der relativen Anordnung wie in Abbildung 4 hängt von der jeweiligen Implementierung ab. Die bekannteste ist Trudgills Gravitationsmodell,⁸ das „a rather simple formula sometimes employed by geographers“ (Trudgill 1974: 233) auf die Geolinguistik anwendet – wenn diese auch nicht näher theoretisch begründet wird (cf. Pickl 2013a: 219).⁹ Die Formel für sprachliche Interaktion zwischen zwei Orten ist in perfekter Analogie zum

Newtonschen Gravitationsmodell formuliert (cf. Nerbonne/Heeringa 2007: 31):

$$M_{ij} = s \cdot \frac{P_i \cdot P_j}{d_{ij}^2} \quad (1)$$

Hier steht M_{ij} für die sprachliche Interaktion zwischen zwei Orten i und j , d_{ij} für die geographische Distanz zwischen i und j , und P_i bzw. P_j für die Population des jeweiligen Ortes. s ist ein Proportionalitätsfaktor, der für „prior-existing linguistic similarity“ (Trudgill 1974: 234) steht und den Trudgill einführt, weil „it appears to be psychologically and linguistically easiest to adopt linguistic features from those dialects or accents that most closely resemble ones [!] own, largely, we can assume, because the adjustments that have to be made are smaller“ (Trudgill 1974: 234).

Auf der Grundlage von Trudgills Gravitationsmodell kann man davon ausgehen, dass sich die linguistische Distanz des Gravitationsmodells (d_{grav}) zwischen zwei Orten proportional zum Quadrat der geographischen Distanz und umgekehrt proportional zum Produkt der Populationszahlen verhält (cf. Nerbonne/Heeringa 2007) – und damit invers zur Interaktion des Modells,¹⁰ so dass man näherungsweise postulieren bzw. folgern kann:¹¹

$$d_{\text{grav}} = \frac{d_{ij}^2}{P_i \cdot P_j} \quad (2)$$

Die so errechneten Distanzen können verwendet werden, um den topologischen Raum des Gravitationsmodells zu (re-)konstruieren. Sie sind jedoch, wie bereits erwähnt, nicht metrisch und lassen sich nicht ohne Weiteres in Form eines zweidimensionalen Ortsnetzes abbilden. Für seine Visualisierung verwenden wir deshalb hier Multidimensionale Skalierung (MDS), eine u. a. in der Dialektometrie häufig verwendete statistische Standardmethode zur Transformation von nicht-metrischen in euklidische Distanzen. Dabei werden die paarweisen Beziehungen zwischen den einzelnen Ortspunkten (so viel wie nötig, so wenig wie möglich) verzerrt, um den Orten unter weitgehender Bewahrung ihrer paarweisen Relationen Koordinaten in einem Ortsnetz zuweisen und ihre relative Anordnung so visualisieren zu können (siehe Abbildung 5). Das Bestimmtheitsmaß R^2 gibt dabei den durch die transformierten Distanzen erfassten Anteil an den originalen (i. e. in diesem Fall topologischen) Relationen an.

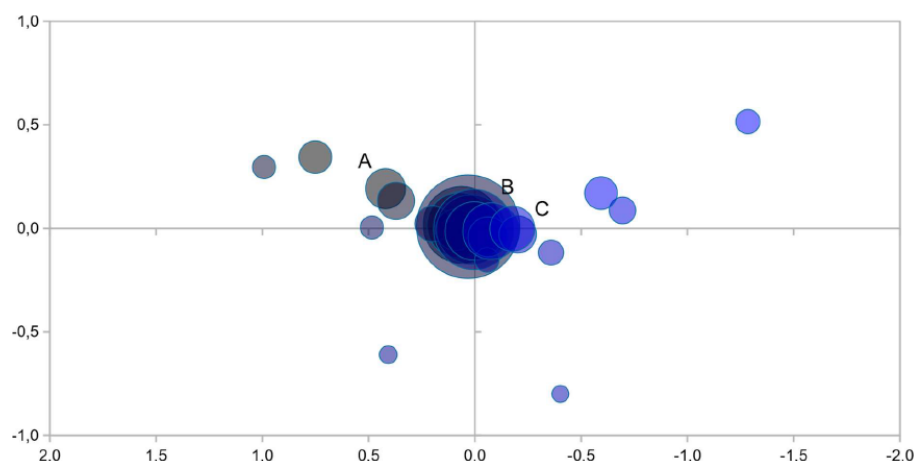


Abbildung 5

Topologische Struktur des Gravitationsmodells (Beispiel). Visualisierung mit MDS ($R^2 = 0,856$)

Die Anordnung der Ortspunkte in Abbildung 5 bewahrt die mit dem Gravitationsmodell errechneten paarweisen Relationen zu 85,6 %, stellt also eine gute Näherung der zugrundeliegenden multidimensionalen, nichteuklidischen Topologie des Modells dar. Der „Gravitationseffekt“ ist (im Vergleich zur Beispieldarstellung in Abbildung 3) erkennbar sehr ausgeprägt. Die größeren Orte sind aufgrund der starken „Anziehungskraft“ nicht nur näher zusammengedrückt, sondern haben sich zu einem praktisch untrennbaren zentralen Cluster vereinigt, zu dem auch die Orte B und C gehören. Lediglich die periphersten der kleineren Orte, darunter Ort A, sind von diesem zentralen Cluster relativ losgelöst, wenn auch klar daran orientiert.

2.3.2 Probabilistisches Modell

Eine weitere Operationalisierung des interaktionalen Ansatzes stellt das probabilistische Modell von Pickl (2013a) dar, das bislang noch nicht anhand von empirischen Daten erprobt wurde. Diesem Ansatz liegt ein umfassenderes interaktionales und probabilistisches Modell von Variation und Wandel im Sprachraum zugrunde, das als interpretativer Rahmen für die quantitative Auswertung von Dialektdaten entwickelt wurde und anders als Trudgills Gravitationsmodell aufgrund expliziter theoretischer Vorannahmen formuliert bzw. postuliert wurde. Das Modell erlaubt unter anderem die Quantifikation „des Kommunikationsanteils oder der ‚Kommunikationswahrscheinlichkeit‘“ (Pickl 2013a: 219) $\#$ eines Orts mit einem anderen. Dieser Wert $\#(a_i, a_j)$ drückt aus, wie wahrscheinlich es ist, dass eine beliebige Person aus Ort a_i mit einer Person aus demselben Ort oder aus Ort a_j interagiert:¹² „Wenn eine Person aus Ort .1 eine Unterhaltung führt, wie groß ist dann die Wahrscheinlichkeit, dass dies mit einer Person aus demselben Ort geschieht, und wie groß, dass dies mit einer Person aus Ort .2 (oder .3, ...) geschieht?“ (Pickl 2013a: 217). Wie bei Trudgills Definition des

Einflusses von einem Ort auf den anderen handelt es sich deshalb um eine asymmetrische Beziehung. Die Kommunikationswahrscheinlichkeit $\kappa^\#$ ist definiert in Anlehnung an die geostatistische Methode der Intensitätsschätzung (cf. Rumpf et al. 2009; Pickl 2013a), wobei durch den Term im Nenner der Formel auch die Bevölkerungsverteilung im Umfeld der fraglichen Orte berücksichtigt wird (Pickl 2013a: 219–220):

$$\tilde{\kappa}(a_i, a_j) = \frac{K(d(a_i, a_j), h) \cdot P(a_j)}{\sum_{a \in A} K(d(a_i, a), h) \cdot P(a)} \quad (3)$$

Die Kommunikationswahrscheinlichkeit von Ort a_i mit Ort a_j , $\kappa^\#(a_i, a_j)$, quantifiziert den erwarteten relativen sprachlichen Input von Ort a_j für Ort a_i . Dabei gibt $P(a)$ die Bevölkerungsgröße eines Ortes an und $d(a_i, a_j)$ die Distanz zwischen zwei Orten.¹³ A ist die Menge aller Orte im jeweils betrachteten Untersuchungsgebiet. $K(d(a_i, a_j), h)$ bezeichnet die sogenannte Kernfunktion, die bestimmt, wie der geographische Abstand zwischen zwei Orten ihre Kommunikationswahrscheinlichkeit – unter Maßgabe eines Glättungsparameters h – beeinflusst. In Pickl (2013a) wird dafür die zweidimensionale (bivariate) Gaußsche Normalverteilung (4) verwendet. Diese Wahl impliziert, dass die Kommunikationswahrscheinlichkeit eines Ortes mit einem anderen mit wachsender Distanz abnimmt, und zwar zunächst nur allmählich, dann aber deutlicher, bis sie in größerer Entfernung nahe Null geht. In welchem Maß die Relevanz eines Ortes für seine Umgebung mit wachsender Entfernung abnimmt, wird durch die Bandbreite h angegeben – je größer dieser Parameter, umso weiter reicht der Einfluss eines Ortes auf seine Umgebung.

$$K_{\text{Gauß}}(d, h) = \frac{1}{2\pi} \cdot e^{-\frac{1}{2} \cdot \left(\frac{d}{h}\right)^2} \quad (4)$$

Durch die Möglichkeit der Wahl der Kernfunktion K und der Bandbreite h ist das Modell einerseits sehr flexibel, andererseits erfordert es entsprechende Vorannahmen bzw. eine Anpassung auf die jeweiligen Daten.

Da das Modell in der in Pickl (2013a) eingeführten Form (3) asymmetrisch ist, ist es zunächst nicht per se topologisch: Es beschreibt keine Distanzen zwischen Orten, sondern gerichtete Beziehungen. Deshalb ist für die hier verfolgten Zwecke eine leichte Anpassung nötig, um auf der Grundlage dieses Modells eine Topologie zu erstellen. Eine naheliegende Methode, um aus den asymmetrischen Kommunikationswahrscheinlichkeiten die symmetrische Kommunikationsdichte, „die Häufigkeit und Intensität

der Kommunikation zwischen zwei sprachlichen Akteuren“ (Pickl 2013a: 220; cf. auch Bloomfield 1933: 46; Schmidt 2010: 213; Pickl 2013a: 54–55), zu ermitteln, ist die Berechnung des geometrischen Mittels der paarweisen Kommunikationswahrscheinlichkeiten:

$$\kappa(a_i, a_j) = \sqrt{\bar{\kappa}(a_i, a_j) \cdot \bar{\kappa}(a_j, a_i)} \quad (5)$$

Ein topologischer Wert lässt sich nun errechnen, indem diese paarweisen Werte in Ähnlichkeiten $\in [0; 1]$ konvertiert werden, indem sie so normiert werden, dass die Ähnlichkeit $\text{prob}(a_i, a_j)$ für jeden Ort mit sich selbst eins 1 ergibt.

$$s_{\text{prob}}(a_i, a_j) = \frac{\kappa(a_i, a_j)}{\sqrt{\kappa(a_i, a_i) \cdot \kappa(a_j, a_j)}} = \frac{K(d(a_i, a_j), h)}{K(0, h)} \quad (6)$$

Diese Ähnlichkeiten lassen sich nun einfach durch Komplementärwertbildung in Distanzen überführen.

Durch die Normierung sind bei s_{prob} und d_{prob} die Gewichtungen durch Bevölkerungszahlen und -verteilung vollständig weggefallen. Dies ist ein nicht erwünschter Effekt, da sich die Basis des Modells so wieder ausschließlich auf die geographische Distanz reduziert. Deshalb wählen wir hier eine Berücksichtigung der Populationszahlen auf andere Weise. Anstatt die Kernfunktionen mit den Bevölkerungszahlen zu gewichten, wird die Bandbreite der Kernfunktion abhängig von der Bevölkerungszahl moduliert. Anschaulich formuliert bedeutet das, dass große Orte eine größere räumliche kommunikative Reichweite erhalten als kleine. Hierzu wird statt einer konstanten Bandbreite h eine variable Bandbreite (cf. Jones 1990) von $h = b \cdot P(a_j)$ mit einem von Fall zu Fall festzulegenden Bandbreitenfaktor b gewählt, da Importar imangenangenommen wird, dass die geographische Reichweite von Orten mit ihrer Bevölkerungsgröße linear zunimmt. Dadurch ergibt sich

$$\bar{\kappa}(a_i, a_j) = \frac{1}{\sum_{a \in A} K(d(a_i, a), b \cdot P(a))} \cdot K_{\text{Gauß}}(d(a_i, a_j), b \cdot P(a_j)) \quad (8)$$

$$d_{\text{prob}}(a_i, a_j) = 1 - s_{\text{prob}}(a_i, a_j) = 1 - \frac{\kappa(a_i, a_j)}{\sqrt{\kappa(a_i, a_i) \cdot \kappa(a_j, a_j)}} = 1 - \sqrt{e^{-\frac{1}{2} \cdot \left(\frac{d(a_i, a_j)}{b \cdot P(a_j)}\right)^2} \cdot e^{-\frac{1}{2} \cdot \left(\frac{d(a_i, a_i)}{b \cdot P(a_i)}\right)^2}} \quad (9)$$

Damit ist die sprachliche Distanz des probabilistischen Modells abhängig von geographischer Distanz und Bevölkerungsverteilung definiert.

Für die Anwendung auf das oben besprochene Beispiel wurde (aufgrund der Eindimensionalität des fiktiven Beispielsraums) der univariate normalverteilte Kerndichteschätzer gewählt und eine variable, von der Einwohnerzahl abhängige Bandbreite (cf. Jones 1990) mit $b = 0,005$ festgelegt (i. e. pro zusätzlichem Einwohner bzw. zusätzlicher Einwohnerin „strahlt“ ein Ort um 0,005 Entfernungseinheiten weiter aus). Durch Einsetzen der jeweiligen Werte in (9) erhält man die paarweisen Distanzen des probabilistischen Modells. Wiederum wurde MDS angewandt, um diese Distanzen in den zweidimensionalen Raum zu projizieren (siehe Abbildung 6).

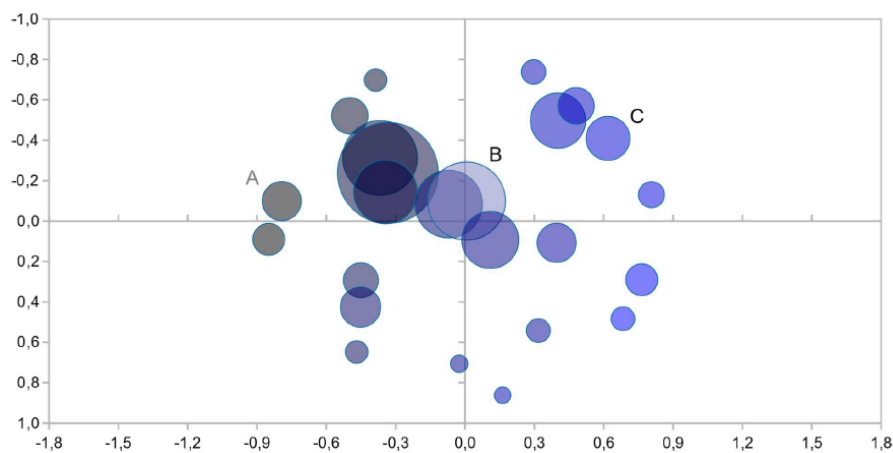


Abbildung 6
Topologische Struktur des probabilistischen Modells
(Beispiel). Visualisierung mit MDS ($R^2 = 0,908$)

Im Vergleich mit Abbildung 5 zeigt sich: Die Distanzen der beiden Formen des probabilistischen Modells lassen sich etwas besser im zweidimensionalen Raum wiedergeben als die des Gravitationsmodells, wobei die Treue des MDS-Modells zu den topologischen Distanzen 90,8 % beträgt. Die Anordnung der Orte lässt in Grundzügen eine Struktur wie in Abbildung 4 oder Abbildung 5 erkennen, der Clustereffekt der bevölkerungsstarken Orte bzw. die Anziehung ist aber geringer als beim Gravitationsmodell; mit der Struktur der Relationen besteht im Detail nur wenig Ähnlichkeit. Die Zentralorte bilden hier ebenfalls einen zentralen Cluster, der jedoch deutlich differenzierter erscheint als beim Gravitationsmodell. Sie treten in eine mittlere (bzw. Vermittler-)Stellung zwischen die kleineren Orte, deren relative Distanzen untereinander weitgehend erhalten bleiben, wodurch ihre ursprünglich horizontale Anordnung um die Zentralorte herum gekrümmt ist.

Welches der Modelle unter Anwendungsgesichtspunkten auf reale Daten ‚realistischer‘ ist, lässt sich aufgrund dieser Beispieldarstellungen freilich nicht beantworten. Der Rest dieses Beitrags ist deshalb der Validierung der verschiedenen Operationalisierungen anhand von empirischen Daten gewidmet. Dabei dient der Zusammenhang zwischen

sprachlicher Ähnlichkeit und Nähe im topologischen Sprachraum als Richtschnur, indem Erstere als Maß für Letztere verwendet wird. Ziel der Modellierung ist es, die reale sprachliche Ähnlichkeit, die aufgrund empirischer Daten bestimmt werden kann, so gut wie möglich unter Zuhilfenahme der verfügbaren Größen Abstand und Bevölkerungsgröße vorherzusagen. Die Qualität des jeweiligen Modells ergibt sich demnach aus dem Grad, zu dem die modellierten Distanzen und die tatsächliche sprachliche Unähnlichkeit bzw. Distanzen übereinstimmen bzw. korrelieren.¹⁴

3 Daten

Datengrundlage für die nachfolgenden Untersuchungen sind die Erhebungen des *Atlas zur deutschen Alltagssprache* (AdA) (Elspaß/Möller 2003–).¹⁵ Die Daten des AdA werden per indirekter Methode in einem Crowdsourcing-Verfahren über Internetfragebögen erhoben. Ziel ist dabei nicht, den individuellen Sprachgebrauch der Probandinnen und Probanden selbst zu ermitteln; stattdessen werden die Teilnehmerinnen und Teilnehmer der Umfragen gebeten, anzugeben, „welchen Ausdruck man in Ihrer Stadt normalerweise hören würde – egal, ob es mehr Mundart oder Hochdeutsch ist“. Erfragt wird – in einer Mischung aus offenen und geschlossenen Fragen – die Gebräuchlichkeit von Varianten vornehmlich lexikalischer, daneben aber auch lautlicher, grammatischer und pragmatischer (z. B. Anrede und Grußformeln) Variablen. Wichtig ist zu betonen, dass es sich um Daten **wahrgenommenen** Sprachgebrauchs handelt. Wie aus der Fragestellung hervorgeht, liegt der Fokus auf Sprachgebrauch im urbanen Bereich. Die Befragungen sind offen, sodass grundsätzlich jede/r teilnehmen kann. Während das Geschlechterverhältnis ungefähr ausgeglichen ist und grundsätzlich alle Altersgruppen vertreten sind, werden durch das Medium der Onlinebefragung hauptsächlich Probandinnen und Probanden der jüngeren und mittleren Generation erreicht (was von den beiden AdA-Bearbeitern durchaus intendiert war): Etwa drei Viertel der bis zu 20 000 Teilnehmerinnen und Teilnehmer in den bisher zwölf Befragungsrunden waren jünger als 40 Jahre alt. Mit „Alltagssprache“ sind im Rahmen des AdA die in der informellen Kommunikation verwendeten Sprachformen gemeint, die sowohl im Nähebereich von Familie und engen Freundinnen und Freunden üblich sind als auch darüber hinaus unter nicht näher miteinander bekannten Sprechenden in einem lokalen Kontext (z. B. in einem örtlichen Café oder Geschäft). Damit ist die namengebende Alltagssprache funktional definiert, nicht formal.

Für die vorliegende Untersuchung wurde auf eine gezielte Auswahl der AdA-Daten zurückgegriffen: Zum Einen wurden nur Daten aus Österreich verwendet, um den Untersuchungsraum und die räumlich-topologischen Relationen im Sinne einer Fallstudie übersichtlich zu halten. Österreich bietet sich aufgrund seines spezifischen Zusammenspiels aus Geographie und Bevölkerungsdichte als Testgebiet

an. Zum Anderen wurden nur die Daten aus den Erhebungsrunden 8–10 verwendet, da die früheren Erhebungsrunden keine ausreichenden geographischen Informationen für diese geolinguistische Untersuchung enthalten – erst in den jüngeren Runden wurden die exakten Postleitzahlen abgefragt, die für eine ausreichende geographische Zuordnung notwendig sind: Die Ortspunkte sind durch Ortsname und Postleitzahl eindeutig identifizierbar und wurden auf dieser Grundlage händisch mit geographischen Koordinaten versehen. Da das Ortsnetz der verschiedenen Runden nicht einheitlich ist (schließlich sind die Probandinnen und Probanden von Runde zu Runde nicht völlig identisch), wurde es homogenisiert, indem es auf diejenigen Orte in Österreich beschränkt wurde, die tatsächlich für alle drei betrachteten Runden Datensätze enthalten. Daraus ergibt sich eine Datengrundlage, die aus 124 Orten, 191 Variablen und 4 114 Varianten besteht und auf 165 268 Einzelantworten durch die Gewährspersonen beruht (cf. auch Pickl et al. 2019: 51).

Die Umfrageergebnisse des AdA liegen in Form von Tabellen vor, in denen die nominalskalierten Antworten der einzelnen Teilnehmerinnen und Teilnehmer unter Angabe der jeweiligen Postleitzahl verzeichnet sind; mit anderen Worten, die Datenbank besteht aus einer Zeile für jede Gewährsperson und einer Spalte für jede abgefragte Variable (für die in Runden 8–10 enthaltenen Variablen und Varianten cf. Elspaß/Möller 2003–). Die einzelnen Zellen enthalten die Antworten der Gewährspersonen (cf. Tabelle 1).

Tabelle 1
Basisdatenstruktur des AdA

informant id	postal code	variable 1	variable 2	
1	86150	<i>variant a</i>	<i>variant b</i>	...
2	86167	<i>variant a</i>	<i>variant a</i>	...
3	86167	<i>variant b</i>	<i>variant c</i>	...
...	...	...	...	...

Das Ortsnetz wird hier zunächst hinsichtlich seiner geographischen Struktur dargestellt. Die einzelnen Ortspunkte sind durch ihre geographischen Koordinaten definiert und mit einem Wert für ihre jeweilige Bevölkerungszahl (Gemeindegröße) versehen (cf. Statistik Austria (2020)). Die Ortspunkte sind nach ihrer geographischen Lage (Koordinaten) und Einwohnerzahl (Blasengröße) in Abbildung 7 dargestellt. Die Farbgebung basiert auf der geographischen Lage der jeweiligen Orte und soll sie in den nachfolgenden Abbildungen tendenziell identifizierbar machen und ihre geographische Einordnung erleichtern.

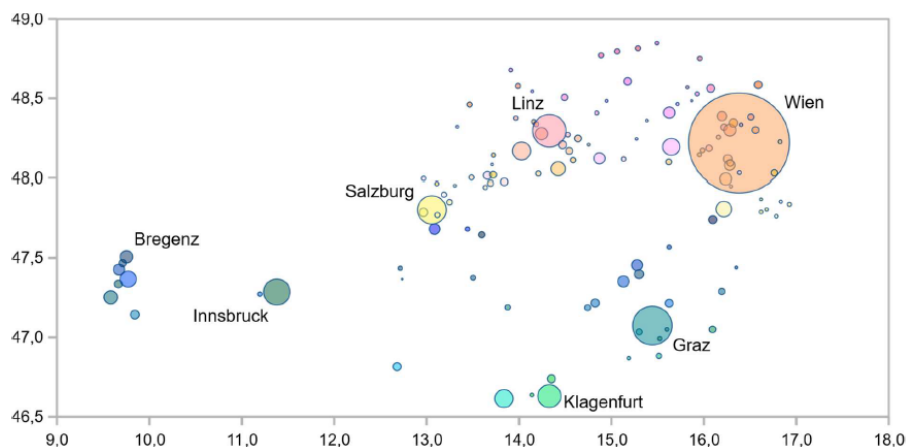


Abbildung 7

Die Ortspunkte im Untersuchungsgebiet nach geographischen Koordinaten (Position) und Einwohnerzahlen (Blasengröße)

Deutlich wird bei dieser Darstellung die Dominanz Wiens mit Blick auf seine schiere Größe – mehr als ein Fünftel der Bevölkerung Österreichs lebt in Wien –, die allerdings mit einer rein geographisch gesehen eher peripheren Lage verbunden ist. Die regionalen Zentren Graz, Linz, Salzburg, Innsbruck, Klagenfurt etc. sind dagegen von deutlich geringerer, aber untereinander durchaus vergleichbarer Größe und über das Staatsgebiet Österreichs mehr oder weniger gleichmäßig verteilt. Mit Blick auf die kleineren Orte ist zu konstatieren, dass die Belegorte in Abbildung 7 ungleichmäßig verteilt sind. Dies spiegelt einerseits die Bevölkerungsverteilung in Österreich wieder, die durch den Alpenhauptkamm mit entsprechend dünn besiedelten Gebieten geprägt ist, andererseits aber auch die wenig gesteuerte Datenerhebung, bei der eine homogene Gebietsabdeckung nicht garantiert ist. Einen Sonderfall stellt das ganz im Westen gelegene Vorarlberg dar, das einerseits das einzige Gebiet mit alemannischem Basisdialekt in Österreich darstellt (siehe Abbildung 10, Abschnitt 4.1) und andererseits im vorliegenden Untersuchungsgebiet durch seine periphere Lage und die hohe Dichte an Orten mittlerer Größe geprägt ist.

4 Ergebnisse

Wie in Abschnitt 2.3.2 erwähnt, dient als Grundlage für die Validierung der Modelle die sprachliche (Un-)Ähnlichkeit zwischen den einzelnen Ortspunkten in Form ihrer linguistischen Abstände. Zur Berechnung der linguistischen Distanzen wurden die sprachlichen Belegdaten in relative Pseudofrequenzen konvertiert, so dass die Summe für alle Varianten einer Variablen am selben Ort 1 ergibt (siehe Tabelle 2; zum Vorgehen cf. Pickl 2013a: 83–85; Pröll/Elspaß/Pickl im Druck).

Tabelle 2

Transformierte Datenstruktur (relative Variantenfrequenz pro Ort)

location	relative frequency of variant a of variable 1	relative frequency of variant b of variable 1	relative frequency of variant a of variable 2	...
<i>location name 1</i>	0.10	0.90	0.05	...
<i>location name 2</i>	0.25	0.75	0.33	...
<i>location name 3</i>	0.00	1.00	0.40	...
...	...	...	...	...

Auf dieser Grundlage wurden zwischen allen Ortspunkten paarweise Manhattan-Distanzen, auch bekannt als City-Block-Distanzen (cf. Speelman/Grondelaers/Geeraerts 2003: 320–322; Pickl 2013a: 99–102; Pickl et al. 2014) errechnet, bei denen die Differenzen der Einzelkoordinaten der Ortspunkte summiert werden,¹⁶ um die sprachlichen Unterschiede zu quantifizieren. Die so ermittelten sprachlichen Distanzen wurden einer MDS unterzogen und in Abbildung 8 visualisiert.

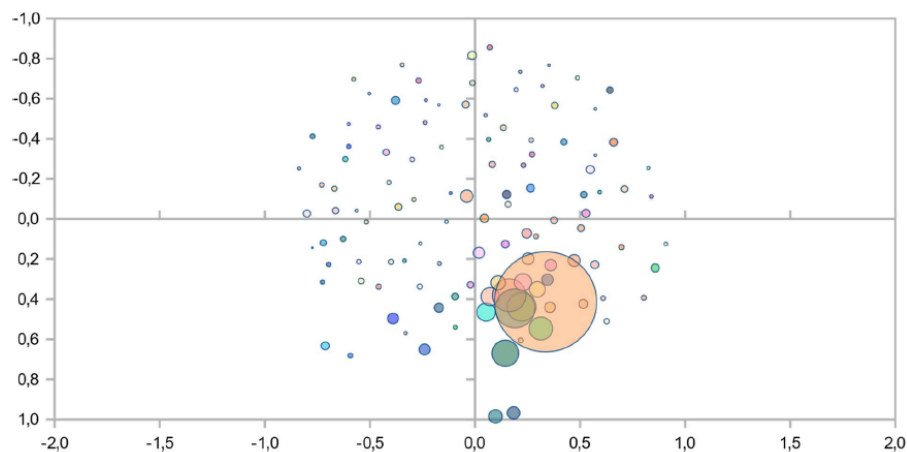


Abbildung 8

MDS-Plot der linguistischen Distanzen ($R^2 = 0,878$)

Deutlich hervor tritt ein Cluster-Effekt bei den bevölkerungsreichen Städten, während eine Relevanz geographischer Nähe auf die topologische Anordnung (die sich in der Bündelung von Orten mit jeweils ähnlichen Farbwerten zeigen würde) nur ansatzweise zu erkennen ist. Dabei handelt es sich zunächst um einen rein visuellen Eindruck; der Zusammenhang zwischen Geographie und topologischer Lage wird im

Folgenden im Rahmen der verschiedenen Modelle näher beleuchtet. Für das weitgehende Fehlen einer klaren geographisch konditionierten

Anordnung in Abbildung 8 gibt es unseres Erachtens zumindest drei mögliche Erklärungsansätze:

1. Möglicherweise macht sich hier einfach bemerkbar, dass empirisch erhobene Daten immer mit statistischen Fluktuationen und Unsicherheiten verbunden sind, die sich als Rauschen oder „noise“ in den Daten – und damit natürlich auch in entsprechenden Auswertungen – bemerkbar machen.
2. Unter Umständen spielen hier komplexe Zusammenhänge eine Rolle, die sich durch abstrahierende und dadurch vereinfachende Modelle wie das Gravitationsmodell – oder überhaupt durch geolinguistische Theoriebildung – kaum erfassen lassen, die jedoch für eine (nur scheinbar) ungeordnete Konfiguration verantwortlich sind.
3. Außerdem ist es denkbar, dass der Alltagssprachliche Bereich im Vergleich zu den Basisdialekten deutlich weniger geographisch konditioniert ist, sodass andere Faktoren wie z. B. Bevölkerungszusammensetzung, wirtschaftliche Prosperität, Erwerbstätigkeit, durchschnittlicher Bildungsgrad oder Einkommen pro Kopf eine größere Rolle spielen als die geographische Lage.

Mit den vorliegenden (Meta-)Daten kann vor allem der letzte Erklärungsansatz leider nicht überprüft werden; es scheint jedoch sinnvoll, solche Faktoren in zukünftige Modelle einfließen zu lassen und ihre Relevanz empirisch zu überprüfen.

4.1 Wellenmodell (geographische Distanzen)

Im Rahmen des Wellenmodells haben, wie oben diskutiert, ausschließlich die geographischen Distanzen einen Einfluss auf die Lagebeziehungen der Orte untereinander. Um die empirische Plausibilität des Wellenmodells zu überprüfen, wird in diesem Abschnitt der Zusammenhang zwischen geographischen und linguistischen Distanzen untersucht. Wie zu erwarten, weisen die im vorausgehenden Abschnitt ermittelten Alltagssprachlichen Distanzen in der Tat einen deutlichen Zusammenhang mit den geographischen Distanzen auf, wie Abbildung 9 zeigt.

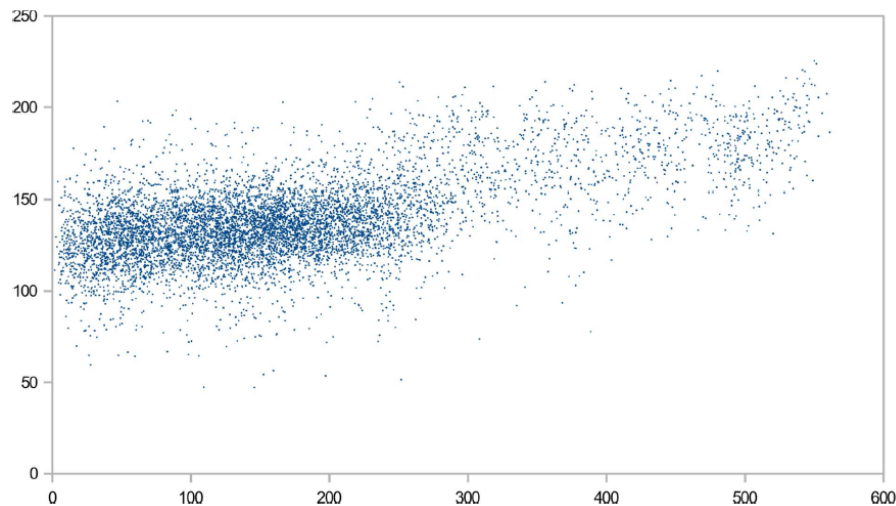


Abbildung 9

Streudiagramm der geographischen Distanzen (in km, x-Achse) und der sprachlichen Distanzen (y-Achse) für alle Ortspaare im Untersuchungsgebiet

Anhand der Verteilung der Punkte, die jeweils für ein Ortspaar stehen, lässt sich erkennen, dass der Zusammenhang von geographischer und linguistischer Distanz insgesamt monoton ist. Dennoch ist etwa ab 300 km ein Bruch zu erkennen, ab dem die linguistischen Distanzen konsistent etwas größer sind als es aufgrund des bisherigen stetigen Verlaufs der Verteilung zu erwarten wäre. Dies entspricht etwa der Entfernung zwischen Wien und den östlichsten Gebieten Tirols. Aufgrund der Struktur des Ortsnetzes repräsentieren alle Diagrammpunkte, die jenseits der 300 km-Marke liegen, Ortspaare, die einen Ort in Tirol oder Vorarlberg und einen im Osten Österreichs umfassen. Der abweichende Verlauf der entsprechenden Verteilung gegenüber der Verteilung unter 300 km deutet also auf eine größere sprachliche Unähnlichkeit hin, die sich nicht allein aufgrund des größeren Abstands erklären lässt und die entsprechende Modelle deshalb vermutlich nicht vorhersagen können; sie scheint auch nicht an traditionelle basisdialektale Grenzen gebunden zu sein, denn der Übergangsbereich zwischen bairischen und alemannischen Gebieten nach Wiesinger (1983) verläuft weiter westlich (siehe Abbildung 10). Dagegen bestätigt eine Faktorenanalyse derselben AdA-Daten in der Tat eine relative alltagssprachliche Eigenständigkeit des Raums Tirol und Vorarlberg gegenüber einem zentralen und einem ostösterreichischen geprägten Raum (cf. Pickl et al. 2019), die sich wohl im erwähnten Bruch in Abbildung 9 niederschlägt.

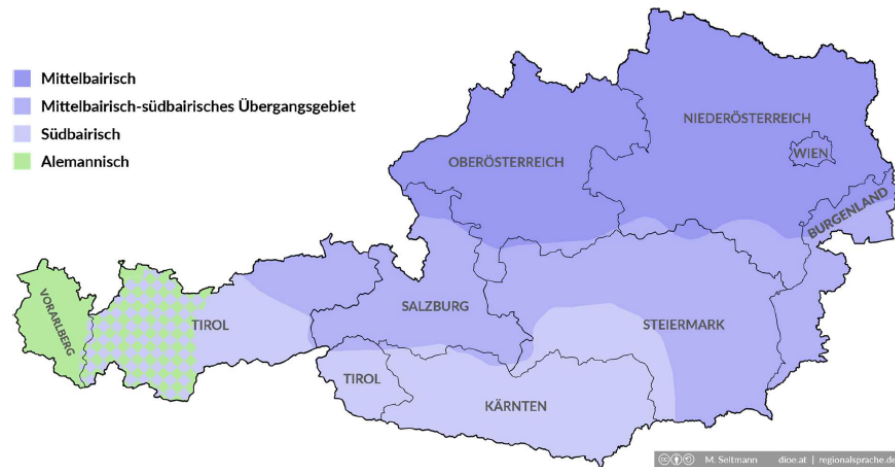


Abbildung 10

Die traditionellen Dialektgebiete in Österreich nach Wiesinger (1983)

Um die Natur des Zusammenhangs zwischen geographischer Distanz und sprachlicher (Un-)Ähnlichkeit möglichst genau zu beleuchten, werden verschiedene Arten möglicher Relationen beleuchtet, indem einzelne Variablen logarithmiert werden (zur gewichteten Korrelation siehe unten). Die lineare Korrelation (ungewichteter Pearsonscher Korrelationskoeffizient) zwischen beiden Größen beträgt 0,595, was einer erklärten Varianz von $r^2 = 0,354$ entspricht (siehe Tabelle 3). Die in anderen Untersuchungen meist beobachteten logarithmischen Zusammenhänge (i. e. zwischen linguistischer Distanz und logarithmierter geographischer Distanz; siehe Abschnitt 2) bestätigen sich hier nicht:¹⁷ Die Korrelation zwischen der logarithmierten geographischen Distanz und der linguistischen Distanz ist mit 0,471 ($r^2 = 0,222$) geringer als der lineare Zusammenhang. Dennoch lässt sich anhand Abbildung 2 leicht erkennen, dass eine logarithmische Relation, vor allem im Kontext der bisherigen Forschung, nicht unplausibel erscheint.¹⁸ Interessanterweise erzielt eine Korrelation zwischen logarithmierter linguistischer Distanz und linearer geographischer Distanz sogar höhere Werte als bei umgekehrter Logarithmierung, die jedoch unter denen der linearen Korrelation liegen.

Tabelle 3

Korrelationen und erklärte Varianzen zwischen geographischer (d_{geo}) und linguistischer Distanz (d_{ling})

	Korrelation (ungewichtet)		Korrelation (gewichtet)	
	r	r^2	r	r^2
linear	0,595	0,354	0,472	0,223
$\log(d_{\text{geo}})$	0,471	0,222	0,368	0,135
$\log(d_{\text{ling}})$	0,553	0,306	0,390	0,152
$\log(d_{\text{geo}}), \log(d_{\text{ling}})$	0,450	0,202	0,304	0,092

Die Ergebnisse der einfachen Pearsonschen Korrelation betreffen die Vorhersagekraft des entsprechenden Regressionsmodells mit Bezug

auf die Alltagssprachliche (Un-)Ähnlichkeit zwischen den einzelnen Ortspunkten. Dabei haben bevölkerungsreiche Orte dasselbe Gewicht für die berechneten Werte wie kleine Orte. Das bedeutet, dass so zwar die Vorhersagequalität mit Bezug auf ganze Orte ermittelt, dadurch aber Einwohnerinnen und Einwohner von größeren Städten systematisch unterrepräsentiert sind. Wenn etwa ein Ort ein vom Modell abweichendes Verhalten aufweist und z. B. systematisch größere Distanzen aufweist als vergleichbare andere, fällt dies für Orte mit einer Bevölkerung von beispielsweise 1 000 genauso ins Gewicht wie für Wien mit seinen 1,9 Millionen Einwohnerinnen und Einwohnern. Aus diesem Grund bietet es sich an, auch jeweils die mit den lokalen Populationen gewichtete Korrelation zu berechnen. Deren Ergebnisse lassen sich dahingehend interpretieren, dass sie die Vorhersagekraft des Modells für die näherungsweise sprachliche (Un-)Ähnlichkeit zwischen allen Einwohnerinnen und Einwohner im Untersuchungsgebiet und nicht zwischen allen Orten im Untersuchungsgebiet quantifizieren. Die entsprechenden Werte sind ebenfalls in Tabelle 3 angegeben. Wenig überraschend schneidet die ungewichtete Korrelation beim Wellenmodell durchweg besser ab als die gewichtete, da sie die Bevölkerungsverteilung unberücksichtigt lässt. Die festgestellte Vorhersagekraft der geographischen Distanzen betrifft also allein die topologischen Relationen zwischen Ortspunkten; bezogen auf einzelne Sprecherinnen und Sprecher schneidet sie deutlich schlechter ab. Vereinfacht lässt sich sagen, dass das Wellenmodell bei der Vorhersage linguistischer Distanzen im ländlichen Raum bessere Vorhersagen trifft als für Städte. Es ist zu erwarten, dass Ansätze, die die Bevölkerungsverteilung in die Modellierung einbeziehen, bessere Werte bei der Vorhersage der sprachlichen (Un-)Ähnlichkeit bezogen auf einzelne Sprecherinnen und Sprecher erzielen.

Die vom Wellenmodell präferierten geographischen Distanzen stellen bei den vorliegenden Daten also einen durchaus brauchbaren Prädiktor für die relationalen sprachlichen Lagebeziehungen der einzelnen Orte dar: Selbst unter Berücksichtigung der Bevölkerungsverteilung bei der Beurteilung der Vorhersagequalität des Modells schneidet der geographische Abstand nicht viel schlechter, bei logarithmischer Regression sogar besser ab als bei einer rein bevölkerungsunabhängigen Berechnung der Korrelation. Dennoch ist einschränkend festzuhalten, dass es offenbar regionale Effekte gibt, die sich in einem gegenüber dem Basistrend erhöhten Abstand zwischen Orten in Tirol und Vorarlberg gegenüber den Orten in den anderen Teilen Österreichs äußern.

4.2 Hierarchisches Modell (Population)

Für das hierarchische Modell sind vor allem Bevölkerungszahlen maßgeblich, räumliche Distanzen spielen allenfalls eine untergeordnete Rolle. Da es keine Operationalisierung des Kaskadenmodells gibt (mit Ausnahme des Gravitationsmodells, das hier unter den Interaktionsmodellen behandelt wird; siehe Abschnitt 2.3) – und zur

besseren Kontrastierung mit dem Wellenmodell einerseits und den interaktionalen Modellen andererseits – wird in diesem Beitrag die Bevölkerungsgröße als einzig relevante Größe für das Kaskadenmodell behandelt. Da die Einwohnerzahl eine ortsbezogene Größe ist, die linguistischen Distanzen *.ling* aber relationale, auf Ortspaare bezogene Werte, verwenden wir als Vergleichsgröße für *.ling* das paarweise geometrische Mittel der Populationszahlen.¹⁹ Auch hier ist ein klarer Zusammenhang mit den linguistischen Distanzen feststellbar (siehe Abbildung 11).

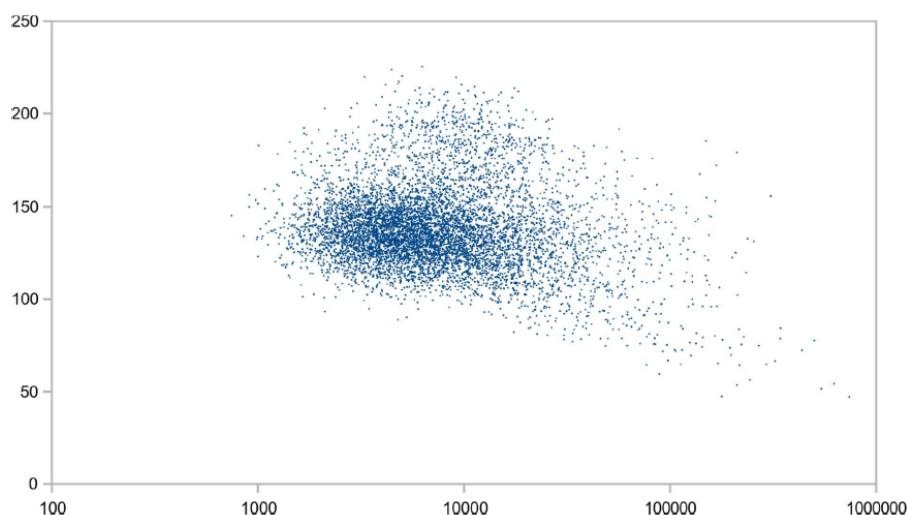


Abbildung 11

Streudiagramm der paarweise geometrisch gemittelten Einwohnerzahlen (x-Achse, logarithmisch²⁰) und der sprachlichen Distanzen (y-Achse) für alle Ortspaare im Untersuchungsgebiet

Wie zu erwarten, sinkt mit steigender Einwohnerzahl die sprachliche Distanz. Das bedeutet, dass Paare aus größeren Orten sprachlich ähnlicher sind als Paare aus kleineren Orten oder als „gemischte“ Paare, die im mittleren Bereich liegen. Dies lässt sich mithilfe der Annahme des Kaskadenmodells erklären, nach der Neuerungen von den größten Orten ausgehend zunächst die nächstgroßen erfassen (wodurch diese stärkere sprachliche Ähnlichkeiten ausbilden) und erst nach und nach zu den kleineren Orten durchsickern, die als letztes erfasst werden und damit am längsten größere Unterschiede zu den größeren Orten sowie zueinander bewahren. Die statistischen Werte für den Zusammenhang zwischen Bevölkerungszahlen und linguistischen Distanzen sind in Tabelle 4 aufgeführt.

Tabelle 4

Korrelationen und erklärte Varianzen zwischen gemittelten
Bevölkerungszahlen ($P\#$) und linguistischer Distanz (d_{ling})

	Korrelation (ungewichtet)		Korrelation (gewichtet)	
	r	r^2	r	r^2
linear	-0,232	0,054	-0,554	0,306
$\log(\bar{P})$	-0,179	0,032	-0,521	0,272
$\log(d_{\text{ling}})$	-0,303	0,092	-0,656	0,431
$\log(\bar{P}), \log(d_{\text{ling}})$	-0,230	0,053	-0,576	0,331

Die ungewichteten Zusammenhänge liegen hier um einiges niedriger als beim Wellenmodell; durch Gewichtung mit Populationszahlen lassen sich aber teils höhere Werte als das Wellenmodell überhaupt erreicht erzielen, v. a. bei logarithmiertem d_{ling} . Dies deutet darauf hin, dass Einwohnerzahlen als Faktor in erster Linie bei Städten (die bei der gewichteten Korrelation stärker ins Gewicht fallen) relevant sind (und dort eine deutlich größere Rolle spielen als die geographische Distanz für kleine Orte), während bei kleineren Orten (die aufgrund ihrer größeren Anzahl bei der ungewichteten Korrelation überrepräsentiert sind) vor allem ihre geographische Lage maßgeblich ist. Mit Blick auf das hierarchische Modell bedeutet das, dass es vor allem für größere Orte und Städte geeignet zu sein scheint, während es mit abnehmender Population an Bedeutung verliert und dort gegenüber dem Wellenmodell an Relevanz einbüßt.

4.3 Interaktionale Modelle

Stellvertretend für den interaktionalen Ansatz werden hier zwei operationalisierte Modelle betrachtet: das Gravitationsmodell nach Trudgill (1974) und das probabilistische Modell nach Pickl (2013a). Das Gravitationsmodell umfasst eine Quantifizierung der (symmetrischen) sprachlichen Interaktion zwischen zwei Orten, die hier verwendet wird (siehe Abschnitt 2.3.1); das probabilistische Modell definiert zunächst nur die (asymmetrische) Kommunikationswahrscheinlichkeit eines Orts mit einem anderen, auf deren Basis für diese Untersuchung eine Definition der (symmetrischen) Kommunikationsdichte formuliert wurde (siehe Abschnitt 2.3.2), die hier Verwendung findet.

4.3.1 Gravitationsmodell

Mit Trudgills Definition der sprachlichen Interaktion zwischen zwei Orten (siehe (1), Abschnitt 2.3.1) lassen sich auf der Grundlage der geographischen Distanzen und der Bevölkerungszahlen Werte für jedes Ortspaar ermitteln, die sich direkt mit den empirischen linguistischen Distanzen vergleichen lassen (siehe Abbildung 12).

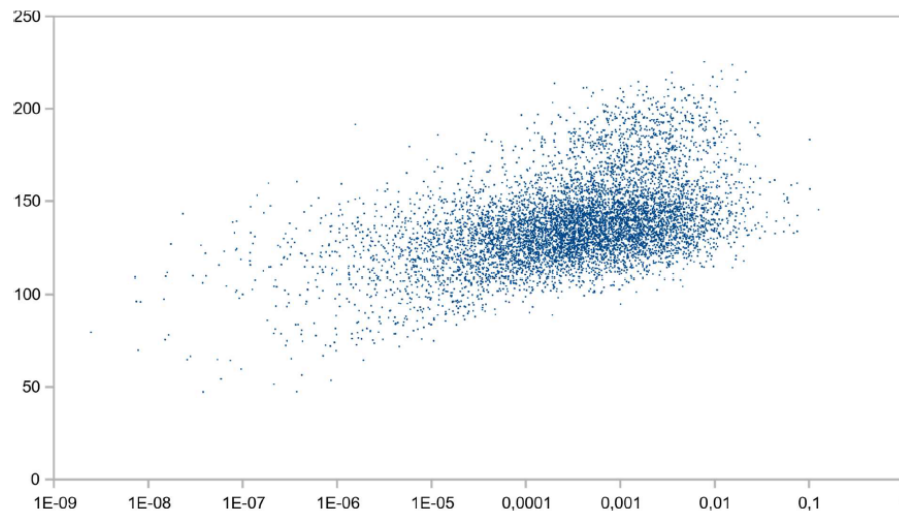


Abbildung 12

Streudiagramm der Interaktionswerte des Gravitationsmodells (x-Achse, logarithmisch) und der sprachlichen Distanzen (y-Achse) für alle Ortspaare im Untersuchungsgebiet

Der visuelle Eindruck weist deutliche Bezüge zu dem des hierarchischen Modells, i. e. der reinen gemittelten geographischen Distanzen (siehe Abbildung 11), auf, wenngleich der Zusammenhang umgekehrter Natur ist (i. e. die Abbildungen tendenziell spiegelverkehrt erscheinen). Dies ist ein möglicher Hinweis darauf, dass das Gravitationsmodell trotz der Berücksichtigung beider Größen dem hierarchischen Modell näher steht und Bevölkerungszahlen in ihm stärker repräsentiert sind als die geographische Lage.

Tabelle 5

Korrelationen und erklärte Varianzen zwischen den Distanzwerten des Gravitationsmodells (d_{grav}) und linguistischer Distanz (d_{ling})

	Korrelation (ungewichtet)		Korrelation (gewichtet)	
	r	r^2	r	r^2
linear	0,193	0,037	0,244	0,060
$\log(d_{\text{grav}})$	0,468	0,219	0,647	0,419
$\log(d_{\text{ling}})$	0,191	0,037	0,217	0,047
$\log(d_{\text{grav}}), \log(d_{\text{ling}})$	0,493	0,243	0,659	0,434

Insgesamt übertrifft das Gravitationsmodell nur leicht den höchsten Wert des hierarchischen Modells (durch Logarithmierung beider Variablen), und zwar bei den gewichteten Korrelationen. So scheint es zunächst keine wesentlich besseren Vorhersagen (zumindest mit Blick auf die unterschiedliche Verteilung der Bevölkerung unter den Orten) zu treffen. Allerdings zeigt sich, dass das Gravitationsmodell bei den ungewichteten Korrelationen (durch Logarithmierung von d_{grav} oder beider Variablen) höhere Werte erzielt als das hierarchische Modell ²¹ und so bessere Vorhersagen auch für kleinere Orte trifft, wenngleich diese Werte nicht an die des Wellenmodells herankommen. Somit lässt

sich durchaus konstatieren, dass durch die Einbeziehung beider Größen – der Bevölkerungszahlen und der geographischen Distanzen – das Gravitationsmodell eine ausgewogenere Modellierung mit Blick auf den ländlichen und den urbanen Raum erzielt.

Da sich das Gravitationsmodell als topologisches Modell verstehen lässt, ist auch eine Visualisierung der sich aus ihm ergebenden Lagebeziehungen naheliegend. Dafür werden die Interaktionswerte für die einzelnen Ortspaare als Ähnlichkeiten einer MDS unterzogen (siehe Abbildung 13).

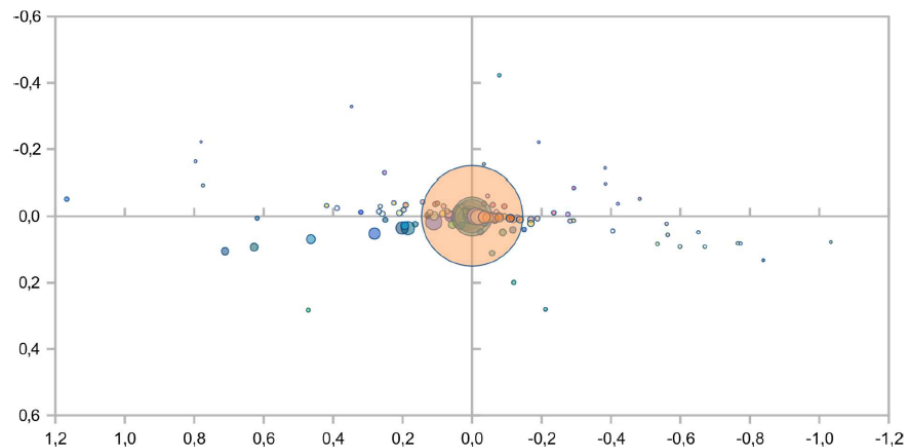


Abbildung 13

MDS-Plot ($R^2 = 0,488$) der Distanzwerte des Gravitationsmodells

Die zweidimensionale MDS bildet die ursprünglichen Relationen des Gravitationsmodells zu 48,8 % ab – ein vergleichsweise niedriger Wert, i. e. die zugrunde liegenden topologischen Relationen sind deutlich komplexer, als im MDS-Plot darstellbar. Zu erkennen ist, wie bereits in Abbildung 5, eine starke Clusterbildung der größeren Orte. Tatsächlich topologisch verstreut finden sich nur die kleinsten Orte sowie die im Ortsnetz peripher gelegenen Orte in Vorarlberg, die sich hier entlang einer Linie zum zentralen Cluster orientieren. Davon abgesehen ist eine Relevanz der geographischen für die topologische Lage visuell kaum festzustellen.

4.3.2 Probabilistisches Modell

Für eine Anwendung des probabilistischen Modells ist es erforderlich, einen Kern und eine Bandbreite zu spezifizieren. Als Kern wird der übliche zweidimensionale Normalverteilungskern verwendet, die Bandbreite wird als variabel in Abhängigkeit von der Bevölkerungsgröße definiert. Damit ist der Bandbreitenfaktor σ (siehe Abschnitt 2.3.2) festzulegen, der hier so gewählt wird, dass die erklärte Varianz σ^2 (unabhängig von der Form des Zusammenhangs, i. e. der Frage, ob und welche Variablen logarithmiert werden) maximiert wird, und zwar getrennt nach ungewichteter und gewichteter Korrelation.

Bei ungewichteter Korrelation kann hier der höchste Wert für .. (0,212) erreicht werden, indem allein die abhängige Variable, der linguistische Abstand, logarithmiert wird (siehe Tabelle 6). Der Bandbreitenfaktor, der sich so ergibt, ist $\hat{b} = 0,012$ km, i. e. pro zusätzlichem Einwohner bzw. zusätzlicher Einwohnerin steigt die geschätzte Ausstrahlung eines Orts um 12 m an. Die übrigen $\hat{b}$ -Werte in Tabelle 6 basieren ebenfalls auf diesem Bandbreitenfaktor. Mit diesem Wert ergibt sich eine Relation zwischen der vom Modell vorhergesagten Distanz und der empirischen linguistischen Distanz wie in Abbildung 14 abgebildet.

Tabelle 6

Korrelationen und erklärte Varianzen zwischen den Distanzwerten des probabilistischen Modells und linguistischer Distanz, optimiert auf möglichst hohe ungewichtete Korrelation (hier bei $\log(d_{\text{ling}})$; $\hat{b} = 0,012$ km)

	Korrelation (ungewichtet)		Korrelation (gewichtet)	
	r	r^2	r	r^2
linear	-0,435	0,190	-0,632	0,400
$\log(d_{\text{prob}})$	-0,303	0,092	-0,287	0,082
$\log(d_{\text{ling}})$	-0,461	0,212	-0,629	0,396
$\log(d_{\text{prob}}), \log(d_{\text{ling}})$	-0,286	0,082	-0,250	0,063

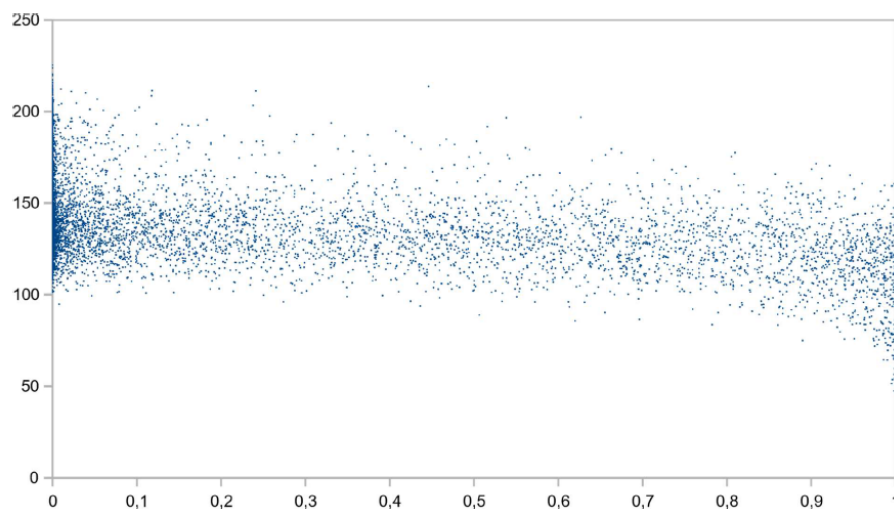


Abbildung 14

Streudiagramm der Distanzwerte des probabilistischen Modells (x-Achse) und der sprachlichen Distanzen (y-Achse) für alle Ortspaare im Untersuchungsgebiet, optimiert auf möglichst hohe ungewichtete Korrelation (siehe Tabelle 6)

Der Zusammenhang stellt sich monoton dar, indem die linguistische Distanz mit wachsender probabilistisch basierter Distanz kontinuierlich abnimmt. Die $\hat{b}$ -Werte mit linearer Kommunikationsdichte sind dabei durchweg höher als bei Logarithmierung (siehe Tabelle 6). Zunächst soll auch hier die Topologie, die sich aus der Anwendung des probabilistischen Modells bei Optimierung auf möglichst hohe

ungewichtete Korrelation (i. e. bei $\lambda = 0,012$ km) ergibt, durch MDS visualisiert werden (siehe Abbildung 15).

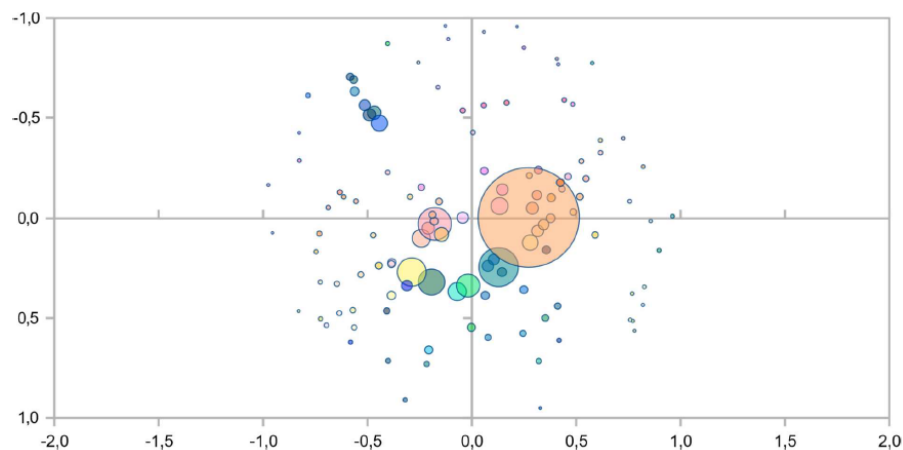


Abbildung 15

MDS-Plot ($R^2 = 0,913$) der Distanzwerte des probabilistischen Modells, gewichtet, optimiert auf möglichst hohe ungewichtete Korrelation (siehe Tabelle 6)

Die ungewichtete Korrelation erreicht hier Werte, die mit denen des Gravitationsmodells ungefähr vergleichbar sind, aber über denen des hierarchischen und unter denen des Wellenmodells liegen. Das bedeutet, dass die bevölkerungsunabhängige Vorhersagekraft, bei der der ländliche Raum stärker repräsentiert ist, bei diesem Modell etwa auf Höhe des Gravitationsmodells liegt und damit schlechter abschneidet als der rein geographische Raum als Prädiktor. Diese Einbuße an regionaler Vorhersagequalität im bevölkerungsunabhängigen Bereich wird quasi „erkauft“ durch eine bessere Repräsentation der Verhältnisse unter Berücksichtigung der Bevölkerungsverteilung, bei der Bewohner und Bewohnerinnen von Städten nicht unterrepräsentiert sind (siehe dazu den nächsten Absatz). Hier zeigt sich wie beim Gravitationsmodell eine Struktur mit Clustereffekt – i. e. bevölkerungsreiche Orte rücken näher zusammen –, bei der jedoch die ursprüngliche geographische Anordnung der größeren Städte noch deutlich erkennbar ist. Sie bilden hier klarer sichtbar einen Zusammenschluss aus Zentralorten, um die herum sich die kleineren Ortschaften gleichsam als Peripherie gruppieren. Auffällig ist hier die Anordnung der Vorarlberger Orte, die größtmäßig im mittleren Bereich liegen und hier eine eigene, vom Zentralcluster unabhängige Gruppe bilden, die sich von kleineren Orten deutlich absondert. Dies lässt sich wohl einerseits durch die auch alltagssprachliche dialektale Eigenständigkeit Vorarlbergs innerhalb Österreichs erklären, andererseits aber wohl auch durch die periphere Lage im hier betrachteten Ortsnetz.

Auffällig bei der Ausprägung der einzelnen λ -Werte für $\lambda = 0,012$ km (siehe Tabelle 6) ist, dass die höheren Werte (i. e. für lineare Kommunikationsdichte) bei der gewichteten Korrelation um einiges höher ausfallen als bei der ungewichteten, obwohl der λ -Wert für einen möglichst hohen gewichteten Zusammenhang optimiert wurde. Dies deutet darauf hin, dass bei einer Optimierung für einen möglich

hohen ungewichteten Zusammenhang mit noch höheren β -Werten zu rechnen ist. So ergibt sich ein wesentlich niedrigerer Bandbreitenfaktor von $\beta = 0,003$ km – und in der Tat erreicht β auf diese Weise für logarithmiertes d_{ling} 0,533 (siehe Tabelle 7), was den höchsten bis hierher erreichten β -Wert unter allen Modellen darstellt.

Tabelle 7

Korrelationen und erklärte Varianzen zwischen den Distanzwerten des probabilistischen Modells und linguistischer Distanz, optimiert auf möglichst hohe ungewichtete Korrelation ($\beta = 0,003$)

	Korrelation (ungewichtet)		Korrelation (gewichtet)	
	r	r^2	r	r^2
linear	–0,336	0,113	–0,672	0,452
$\log(d_{\text{grav}})$	–0,281	0,079	–0,320	0,102
$\log(d_{\text{ling}})$	–0,388	0,150	–0,730	0,533
$\log(d_{\text{grav}}), \log(d_{\text{ling}})$	–0,287	0,082	–0,312	0,098

Dies macht deutlich, dass das probabilistische Modell besonders gut zur Vorhersage von Distanzen unter Berücksichtigung der Bevölkerungsgröße bzw. zur Beschreibung der Verhältnisse unter großen Orten geeignet ist; unter Anpassung von β lassen sich jedoch auch für kleine Orte Relationen modellieren, die mit den Resultaten des Gravitationsmodells vergleichbar sind, aber unter den Werten des reinen Wellenmodells liegen. Für regionale Variation ist also nach wie vor das Wellenmodell am geeignetsten, für urbane Räume stellt das probabilistische Modell einen brauchbareren Ansatz dar.

Abbildung 16 zeigt den Zusammenhang zwischen probabilistischer und linguistischer Distanz bei auf gewichtete Regression optimiertem Bandbreitenfaktor ($\beta = 0,003$).

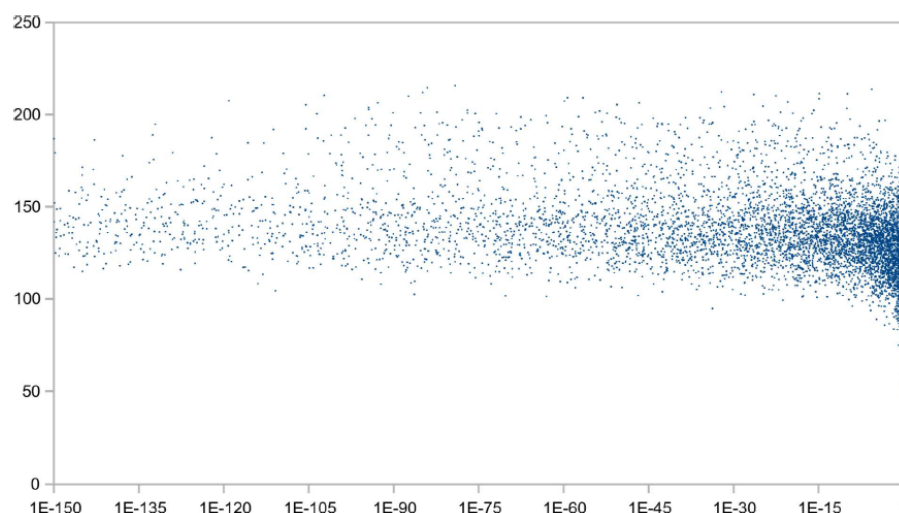


Abbildung 16

Streudiagramm der Distanzwerte des probabilistischen Modells (x-Achse, logarithmisch) und der sprachlichen Distanzen (y-Achse) für alle Ortspaare im Untersuchungsgebiet, optimiert auf möglichst hohe gewichtete Korrelation (siehe Tabelle 7)

Der entsprechende MDS-Plot ist in Abbildung 17 wiedergegeben. Die größeren Orte erscheinen hier in einer periphereren Lage, und ihre Anordnung untereinander entspricht weniger den geographischen Relationen. Visuell besteht hier durchaus eine gewissen Ähnlichkeit zu Abbildung 8, wenngleich Unterschiede u. a. bei der Anordnung der Vorarlberger Orte bestehen.

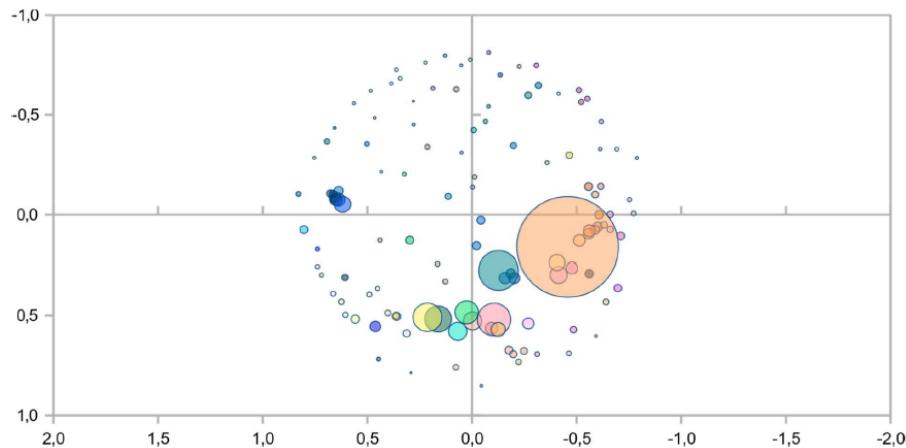


Abbildung 17

MDS-Plot ($R^2 = 0,857$) der Distanzwerte des probabilistischen Modells, gewichtet, optimiert auf möglichst hohe gewichtete Korrelation (siehe Tabelle 7)

4.4 Empirisches Modell

Schließlich soll ausgehend von den empirischen linguistischen Distanzen ein Regressionsmodell auf der Basis von geographischen Distanzen und Bevölkerungsgröße erstellt werden, um das Vorhersagepotenzial dieser beiden Größen auf die sprachlichen Relationen noch gründlicher auszuleuchten. Wie sich bereits an verschiedenen Stellen gezeigt hat, kommt man bei der Modellierung des topologischen Sprachraums nicht umhin, zwei Komponenten zu betrachten, eine regionale und eine urbane. Die regionale Komponente ist durch die geographische Lage gegeben und stellt die horizontale Dimension eines jeglichen Modells dar (durch das Wellenmodell verkörpert), während die urbane Komponente durch die Bevölkerungszahlen spezifiziert ist und die vertikale Dimension darstellt (durch das hierarchische Modell in Reinform repräsentiert). Die verschiedenen bis hierher besprochenen Modelle legen verschiedene Gewichte auf die beiden Komponenten und unterscheiden sich in ihrer Eignung für die Beschreibung der Variation in der jeweiligen Dimension; es konnte jedoch bislang kein Modell identifiziert werden, das für beide Komponenten gleichermaßen ideal ist. Ein Regressionsmodell stellt eine gute Möglichkeit dar, die beiden Komponenten flexibel in die Modellierung zu integrieren und ihren jeweiligen Einfluss auf die abhängige Variable (die linguistische Distanz zu bestimmen).

Ein einfaches lineares Regressionsmodell auf dieser Basis lautet demnach wie folgt:

$$d_{\text{ling}}(a_i, a_j) = \alpha \cdot d_{\text{geo}}(a_i, a_j) + \beta \cdot \bar{P}(a_i, a_j) + \gamma \quad (10)$$

Hier wird die sprachliche Distanz .ling auf der Grundlage der geographischen Distanz .geo und dem geometrischen Mittel der Bevölkerungsgröße $\bar{P}$ modelliert, wobei die Parameter α , β und γ der Anpassung des Modells an die empirischen Daten dienen.²² Dabei können die einzelnen Variablen zur potenziellen Verbesserung der Anpassung logarithmiert werden. Die r - und r^2 -Werte für die einzelnen Anpassungen der einzelnen Regressionsmodelle sind in Tabelle 8 wiedergegeben.

Tabelle 8

Korrelationen und erklärte Varianzen für die verschiedenen Varianten des Regressionsmodell

	Regression (ungewichtet)		Regression (gewichtet)	
	r	r^2	r	r^2
$d_{\text{ling}} = \alpha \cdot d_{\text{geo}} + \beta \cdot \bar{P} + \gamma$	0,652	0,425	0,745	0,555
$d_{\text{ling}} = \alpha \cdot \log(d_{\text{geo}}) + \beta \cdot \bar{P} + \gamma$	0,532	0,283	0,679	0,461
$d_{\text{ling}} = \alpha \cdot d_{\text{geo}} + \beta \cdot \log(\bar{P}) + \gamma$	0,652	0,425	0,727	0,529
$d_{\text{ling}} = \alpha \cdot \log(d_{\text{geo}}) + \beta \cdot \log(\bar{P}) + \gamma$	0,520	0,270	0,648	0,421
$\log(d_{\text{ling}}) = \alpha \cdot d_{\text{geo}} + \beta \cdot \bar{P} + \gamma$	0,647	0,418	0,780	0,608
$\log(d_{\text{ling}}) = \alpha \cdot \log(d_{\text{geo}}) + \beta \cdot \bar{P} + \gamma$	0,550	0,303	0,737	0,543
$\log(d_{\text{ling}}) = \alpha \cdot d_{\text{geo}} + \beta \cdot \log(\bar{P}) + \gamma$	0,635	0,403	0,717	0,515
$\log(d_{\text{ling}}) = \alpha \cdot \log(d_{\text{geo}}) + \beta \cdot \log(\bar{P}) + \gamma$	0,524	0,275	0,660	0,436

Mit Regressionsmodellen lassen sich augenscheinlich sehr gute Anpassungen erzielen. Anders als in Tabelle 6 und Tabelle 7, wo jeweils die gesamte Tabelle für eine Optimierung steht, ist hier jeder einzelne r - bzw. r^2 -Wert das Ergebnis einer individuellen Anpassung, was die im Vergleich insgesamt hohen Werte teilweise erklärt. Die höchsten Werte finden sich jeweils für diejenigen Modelle, bei denen die beiden unabhängigen Variablen linear berücksichtigt wurden. Wie das probabilistische Modell schneiden Regressionsmodelle bei Gewichtung mit den Populationszahlen besser ab als ohne Gewichtung. Insgesamt erzielt die gewichtete Regression bei logarithmierter linguistischer Distanz mit $r^2 = 0,608$ den höchsten Wert, sowohl unter den Regressionsals auch allen anderen hier betrachteten Modellen. Dies ist zweifellos ein Ergebnis der Flexibilität des Modells und der Möglichkeit der Anpassung auf die empirischen Daten. Erstere ist bei den übrigen Modellen überhaupt nicht gegeben, Letztere ausschließlich beim probabilistischen Modell. Ein Nachteil der empirischen Regressionsmodelle ist jedoch die fehlende theoretische Fundierung und damit auch die schlechtere Interpretierbarkeit.

Im Folgenden sollen die Ergebnisse der Regressionsanalysen wie für die übrigen Modelle visualisiert und näher betrachtet werden.

Hierfür wurden stellvertretend zwei Regressionen ausgewählt: diejenige mit den jeweils höchsten r^2 -Werten bei der ungewichteten und der gewichteten Regression. Abbildung 18 zeigt den Zusammenhang zwischen den empirischen dling-Werten und den modellierten Werten der ungewichteten linearen Regression ($r^2 = 0,425$).

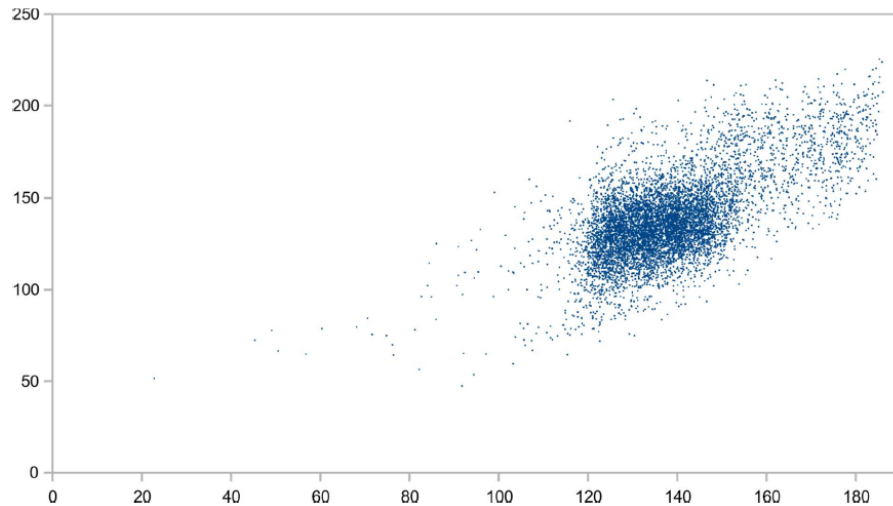


Abbildung 18

Streudiagramm der durch die ungewichteten linearen Regression modellierten Werte (x-Achse) und der sprachlichen Distanzen (y-Achse) für alle Ortspaare im Untersuchungsgebiet

Der entsprechende MDS-Plot (Abbildung 19) zeigt eine Struktur, die weniger als die Topologien des Gravitations- und des probabilistischen Modells von der Anziehungskraft unter den größeren Städten geprägt zu sein scheint. Die geographische Lage der Orte (siehe Abbildung 7) ist bis zu einem bestimmten Grad wiedererkennbar, wenngleich die Distanzen im Detail überformt sind.

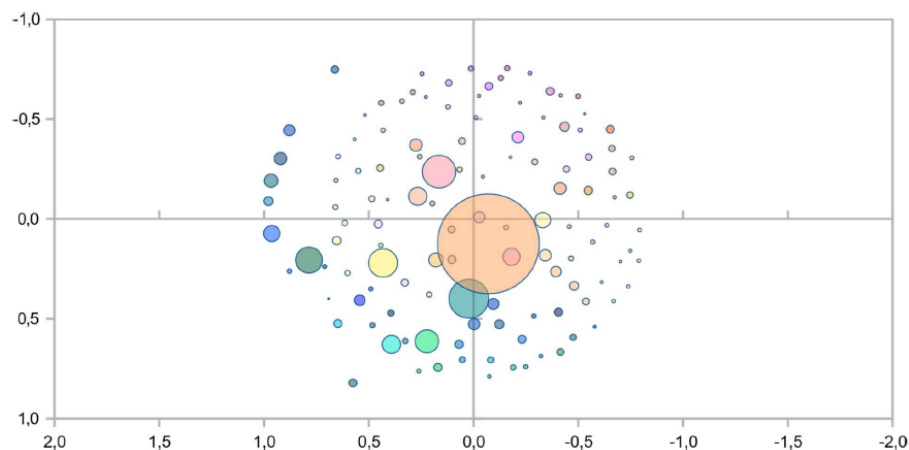


Abbildung 19

MDS-Plot ($R^2 = 0,866$) der durch die ungewichtete lineare Regression modellierten Werte

Abbildung 20 stellt den Zusammenhang zwischen den Werten, die mit der gewichteten Regression mittels logarithmiertem .ling ($r^2 = 0,608$) modelliert wurden, und der linguistischen Distanz dar. Die Unterschiede zu Abbildung 18 sind augenscheinlich gering bzw. liegen im Detail. Dies

liegt zweifellos daran, dass sich die beiden Regressionen mit linearen unabhängigen Variablen kaum unterscheiden. Folglich ist auch der MDS-Plot für die entsprechenden modellierten Werte (Abbildung 21) Abbildung 19 sehr ähnlich.

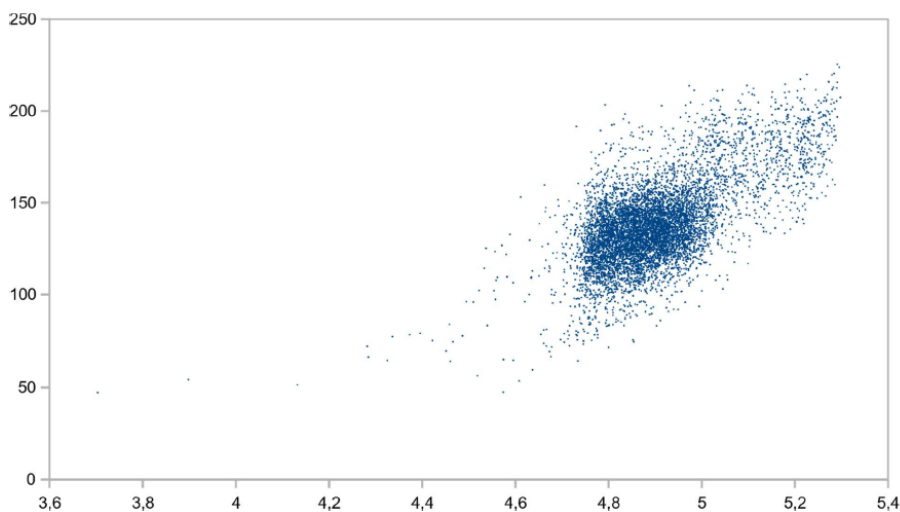


Abbildung 20

Streudiagramm der durch die gewichtete Regression mit logarithmierter linguistischer Distanz modellierten Werte (x-Achse) und der sprachlichen Distanzen (y-Achse) für alle Ortspaare im Untersuchungsgebiet

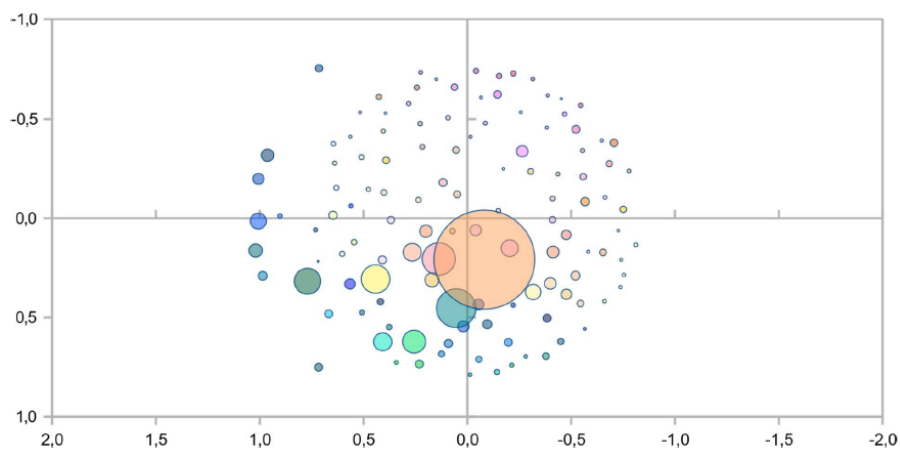


Abbildung 21

MDS-Plot ($R^2 = 0,874$) der durch die gewichtete Regression mit logarithmierter linguistischer Distanz modellierten Werte

5 Diskussion

In den vorausgehenden Abschnitten wurden verschiedene Ansätze, die Topologie des Sprachraums zu modellieren, gegenübergestellt. Dabei hat sich vor allem gezeigt, dass keines der Modelle für sich beanspruchen kann, die Verhältnisse im Sprachraum auf eine objektivierbare Weise besser abzubilden als andere. Das Bild ist wesentlich differenzierter.

Vor allem spielt dabei eine Rolle, dass keines der Modelle regionale (i. e. v. a. geographisch bedingte) **und** interurbane (i. e. v. a.

bevölkerungsabhängige) Komponente der Sprachvariation gleich gut zu beschreiben vermag. Gleichzeitig können allein die Regressionsmodelle die regionale Variation (ungeachtet der Bevölkerungsverteilung, i. e. bei ungewichteter Korrelation) besser erklären als das Wellenmodell. Die interurbane Variation (unter Berücksichtigung der Bevölkerungsverteilung, i. e. bei gewichteter Korrelation) können die meisten Modelle jedoch besser erklären als das Wellenmodell. Letzteres ist also nur dann gut geeignet, wenn man an Unterschieden zwischen Stadt und Land nicht interessiert ist; sobald Dynamiken in den Blick geraten, die von der Bevölkerungsverteilung bzw. dem Urbanitätsgrad abhängen, sind komplexere Modelle erforderlich. Dies zeigt zum einen, dass die Berücksichtigung der Bevölkerungsverteilung nicht nur für die Modellierung selbst, sondern auch für die Bewertung der Qualität der Modelle eine wichtige Rolle spielen muss – ein Umstand, der in bisherigen Studien keine Berücksichtigung fand, in denen stets ausschließlich *..*-Werte aus ungewichteten Regressionen betrachtet wurden, wodurch Bewohner und Bewohnerinnen aus kleinen Orten systematisch überund solche aus großen Orten systematisch unterrepräsentiert waren. Zum anderen wird so deutlich, dass das Verhältnis von städtischem und ländlichem Sprachgebrauch und die sich daraus ergebenden topologischen Strukturen erheblich davon abhängig ist, ob eher die regionale oder die interurbane Komponente der Variation im Vordergrund der Betrachtung steht. Bei der Wahl eines geeigneten Modells – und damit der jeweiligen theoretischen Konzeptualisierung des Sprachraums – spielt u. a. die Abwägung eine Rolle, ob die Verhältnisse im ländlichen Raum oder die unter den Städten eine stärkere Repräsentation finden sollen. Dies zeigt sich z. B. beim probabilistischen Ansatz, bei dem die Wahl des Bandbreitenfaktors *.* darüber entscheidet, ob das Modell eher die Variation unter den größeren Orten präferiert oder eher auf ein ausgewogenes Verhältnis zwischen regionaler und interurbaner Variation abzielt.

Eine ähnliche Abwägung gilt es mit Blick auf die theoretische Fundierung der Modelle zu treffen. Während das Wellen-, das hierarchische und das probabilistische Modell im weitesten Sinne deduktiv, theoretisch fundiert, sind und so eine sinnvolle Interpretation der Ergebnisse erlauben, ist das Gravitationsmodell eher als ein pragmatischer bzw. heuristischer Ansatz zu werten, da die Formulierung der zugrundeliegenden Gleichung nicht auf Deduktion und auch nicht auf Induktion – allenfalls auf Erfahrungswerten aus der Humangeographie – beruht. Die Regressionsmodelle hingegen sind rein induktiv und in diesem Sinne am wenigsten einer sinnvollen Interpretation zugänglich. Ihre große Flexibilität führt zwar zu guten Vorhersagewerten; gleichzeitig haftet ihnen jedoch eine gewisse Beliebigkeit an, da zu viele der Parameter anhand der Daten zu bestimmen sind. Einen Mittelweg stellt das probabilistische Modell dar, das deduktiv hergeleitet und damit theoretisch fundiert ist, eine gewisse Flexibilität durch die Bestimmung eines einzigen Parameters (im Gegensatz zu drei Größen bei der Regression – zusätzlich

zur Wahl zwischen verschiedenen Kombinationen aus linearen und logarithmischen Zusammenhängen) aufweist und gleichzeitig unter allen theoretisch begründeten Modellen die höchsten Vorhersagewerte erzielt (zumindest mit Blick auf die gewichtete Regression).

5 Fazit

In summa bleibt festzuhalten, dass bei der topologischen Modellierung des Sprachraums – und damit auch bei der rein ideellen Konzeptualisierung des Sprachraums – nicht auf die Unterscheidung zwischen einer horizontalen, geographisch bedingten Komponente und einer vertikalen, bevölkerungsabhängigen Komponente verzichtet werden kann. Im Kern läuft das auf eine dreidimensionale Konzeptualisierung des Sprachraums hinaus – bestehend aus dem zweidimensionalen geographischen Raum und einer vertikalen, von der Bevölkerungsverteilung abhängigen Dimension –, in der sowohl die räumlichen Lagebeziehungen der Orte untereinander als auch ihre Populationsgröße eine Rolle spielen. Bei der Beschreibung von Diffusionsvorgängen ist folglich beides ausschlaggebend: die horizontale Diffusion entlang der geographischen Ausdehnung ebenso wie die vertikale Diffusion in Abhängigkeit von der Bevölkerungsverteilung. Dabei zeigt sich, dass der horizontale Diffusionstyp v. a. unter kleineren Orten eine Rolle spielt, während der vertikale Diffusionstyp mit zunehmender Populationsgröße zunehmend an Bedeutung gewinnt.²³ Das heißt, dass hier zwischen Stadt und Land nicht nur ein quantitativer, sondern ein qualitativer Unterschied besteht: Unter großen Städten verbreiten sich Neuerungen anders als in ländlichen Gebieten; im ersten Fall ist der vertikale Diffusionstyp maßgeblich, im zweiten der horizontale. Im Wesentlichen bedeutet dies, dass die Geolinguistik hier nicht ohne zwei Diffusionskonzepte auskommt und deshalb auch verschiedene topologische Konzeptualisierungen der beteiligten Dimensionen benötigt. Während die horizontalen Dimensionen des geographischen Raums symmetrisch sind, ist die vertikale Dimension der Urbanität asymmetrisch und gerichtet. Je größer ein Ort ist, umso mehr ist seine topologische Position von dieser vertikalen Dimension bestimmt; je kleiner er ist, umso mehr verliert diese vertikale Dimension zugunsten der horizontalen Beziehungen an Bedeutung.

Literaturverzeichnis

- Bach, Adolf (1950): *Deutsche Mundartforschung. Ihre Wege, Ergebnisse und Aufgaben. Mit 58 Karten.* 2. Auflage. Heidelberg: Winter. (= *Germanische Bibliothek, Dritte Reihe: Untersuchungen und Einzeldarstellungen*).
- Bailey, Guy et al. (1993): „Some patterns of linguistic diffusion“. *Language Variation and Change* 5: 359–390.
- Becker, Horst (1942): „Über Trichterwirkung, eine besondere Art von Sprachströmung“. *Zeitschrift für Mundartforschung* 18: 59–67.
- Bloomfield, Leonard (1933): *Language*. New York: Holt.

- Boberg, Charles (2000): „Geolinguistic diffusion and the U. S.–Canada border“. *Language Variation and Change* 12: 1–24.
- Britain, David (2002): „Space and spatial diffusion“. In: Chambers, Jack/Trudgill, Peter/Schilling-Estes, Natalie (eds.): *The Handbook of Variation and Change*. Oxford, Blackwell: 603–637.
- Britain, David (2004): „Geolinguistics and linguistic diffusion“. In: Ammon, Ulrich et al. (eds.): *Sociolinguistics. International Handbook of the Science of Language and Society*. Berlin, Mouton de Gruyter: 34–38.
- Britain, David (2010): „Language and space. The variationist approach“. In: Auer, Peter Auer/Schmidt, Jürgen Erich (eds.): *Language and Space. An International Handbook of Linguistic Variation*. Band 1: *Theories and Methods*. Oxford, Blackwell: 142–162.
- Britain, David (2012a): „Countering the urbanist agenda in variationist sociolinguistics: dialect contact, demographic change and the rural-urban dichotomy“. In: Hansen, Sandra et al. (eds.): *Dialectological and Folk Dialectological Concepts of Space. Current Methods and Perspectives in Sociolinguistic Research on Dialect Change*. Berlin/Boston, de Gruyter: 12–30. (= *linguae & litterae* 17).
- Britain, David (2012b): „Innovation diffusion in sociohistorical linguistics“. In: HernándezCampoy, Juan Manuel/Conde-Silvestre, Juan Camilo (eds.): *The Handbook of Historical Sociolinguistics*. Oxford, Blackwell: 451–464.
- Britain, David (2016): „Sedentarism and nomadism in the sociolinguistics of dialect“. In: Coupland, Nikolas (ed.): *Sociolinguistics. Theoretical Debates*. Cambridge, Cambridge University Press: 217–241.
- Chambers, J. K./Trudgill, Peter (1998): *Dialectology*. Cambridge, Cambridge University Press. (= *Cambridge Textbooks in Linguistics*).
- de Lange, Norbert/Nipper, Josef (2018): *Quantitative Methodik in der Geographie*. Paderborn: Schöningh. (= UTB 4933).
- Denis, Derek (2009): *Transmission and Diffusion above the Level of Phonology. Evidence from Thunder Bay*. Masterarbeit, University of Toronto.
- Elspaß, Stephan/Möller, Robert (2003–): *Atlas zur deutschen Alltagssprache* (AdA). atlasalltagssprache.de [22.12.2020].
- Frings, Theodor (1956): *Sprache und Geschichte*. Band 2. Halle: Niemeyer. (= *Mitteldeutsche Studien* 17).
- Goebel, Hans (1984): *Dialektometrische Studien. Anhand italo-romanischer, rätoromanischer und galloromanischer Sprachmaterialien aus AIS und ALF*. Unter Mitarbeit von Siegfried Selberherr, Wolf-Dieter Rase und Hilmar Pudlatz. 3 Bände. Tübingen: Niemeyer.
- Gooskens, Charlotte (2004): „Norwegian dialect distances geographically explained“. In: Gunnarsson, Britt-Louise et al. (eds.): *Language Variation in Europe. Papers from the Second International Conference on Language Variation in Europe ICLAVE 2, June 12–14, 2003*. Uppsala, Uppsala University: 195–206.
- Haggett, Peter (1965). *Locational analysis in human geography*. London: Arnold.
- Hard, Gerhard (2008): „Der *Spatial Turn*, von der Geographie her betrachtet“. In: Döring, Jörg/Thielmann, Tristan (eds.): *Spatial Turn. Das Raumparadigma in den Kultur- und Sozialwissenschaften*. Bielefeld, Transcript: 263–315.

- Heeringa, Wilbert/Nerbonne, John (2001): „Dialect areas and dialect continua“. In: *Language Variation and Change* 13: 375–400.
- Hernández-Campoy, Juan Manuel (2003): „Exposure to contact and the geographical adoption of standard features. Two complementary approaches“. *Language in Society* 32/2: 227–255.
- Kerswill, Paul (1993): „Rural dialect speakers in an urban speech community. The role of dialect contact in defining a sociolinguistic concept“. *International Journal of Applied Linguistics* 3/1: 33–56.
- Labov, William (2003): „Pursuing the cascade model“. In: Britain, David/Cheshire, Jenny (eds.): *Social dialectology: In honour of Peter Trudgill*. Amsterdam, Benjamins: 9–22.
- Labov, William (2007): „Transmission and diffusion“. *Language* 83/2: 344–387.
- Löw, Martina (2001): *Raumsoziologie*. Paderborn: Suhrkamp. (= *Suhrkamp-Taschenbuch Wissenschaft* 1506).
- Möller, Robert/Elspaß, Stephan (2015): „Atlas zur deutschen Alltagssprache“. In: Kehrein, Roland/Lameli, Alfred/Rabanus, Stefan (eds.): *Regionale Variation des Deutschen. Projekte und Perspektiven*. Berlin/Boston, de Gruyter: 519–540.
- Nerbonne, John (2010): „Measuring the diffusion of linguistic change“. *Philosophical Transactions of the Royal Society B: Biological Sciences* 365/1559: 3821–3828.
- Nerbonne, John/Heeringa, Wilbert (2007): „Geographic distributions of linguistic variation reflect dynamics of differentiation“. In: Featherston, Sam/Sternefeld, Wolfgang (eds.): *Roots. Linguistics in Search of its Evidential Base*. Berlin/New York, de Gruyter Mouton: 267–297.
- Nerbonne, John/Kleiweg, Peter (2007): „Toward a dialectological yardstick“. *Journal of Quantitative Linguistics* 14: 148–166.
- Newerkla, Stefan Michael (2013): „Linguistic consequences of Slavic migration to Vienna in the 19th and 20th centuries“. In: Moser, Michael/Polinsky, Maria (eds.): *Slavic Languages in Migration*. Berlin/Wien, LIT Verlag: 247–260. (= *Slavische Sprachgeschichte* 6).
- Olsson, Gunnar (1965): *Distance and Human Interaction*. Philadelphia: Regional Science Research Institute.
- Pickl, Simon (2013a): *Probabilistische Geolinguistik. Geostatistische Analysen lexikalischer Variation in Bayerisch-Schwaben*. Stuttgart: Steiner. (= Zeitschrift für Dialektologie und Linguistik, Beihefte 156).
- Pickl, Simon (2013b): „Verdichtungen im sprachgeografischen Kontinuum“. *Zeitschrift für Dialektologie und Linguistik* 80/1: 1–35.
- Pickl, Simon (2017): „Wann ist eine Grenze eine Grenze? Zur theoretischen Fundierung von Dialektgrenzen und ihrer statistischen Validierung“. In: Christen, Helen/Gilles, Peter/Purschke, Christoph (eds.): *Räume, Grenzen, Übergänge. Akten des 5. Kongresses der Internationalen Gesellschaft für Dialektologie des Deutschen (IGDD)*. Stuttgart, Steiner: 259–281, 402. (= Zeitschrift für Dialektologie und Linguistik, Beihefte 171).
- Pickl, Simon/Elspaß, Stephan (2019): „Historische Soziolinguistik der Stadtsprachen“. In: Pickl, Simon/Elspaß, Stephan (eds.): *Historische*

Soziolinguistik der Stadtsprachen. Kontakt – Variation – Wandel. Heidelberg, Winter: 1–29. (= *Germanistische Bibliothek* 67).

- Pickl, Simon/Pröll, Simon (2019): „Ergebnisse geostatistischer Analysen arealsprachlicher Variation im Deutschen.“ In HSK 30.4: Herrgen, Joachim/Schmidt, Jürgen Erich (eds.): *Sprache und Raum. Ein internationales Handbuch der Sprachvariation*. Teilband 4: *Deutsch*. Unter Mitarbeit von Hanna Fischer und Brigitte Ganswindt. Berlin/Boston, de Gruyter: 861–879. (= *Handbücher zur Sprachund Kommunikationswissenschaft* 30.4).
- Pickl, Simon et al. (2019): „Räumliche Strukturen alltagssprachlicher Variation in Österreich anhand von Daten des ‚Atlas zur deutschen Alltagssprache (AdA)‘“. In: Bülow, Lars/Fischer, Ann-Kathrin/Herbert, Kristina (eds.): *Dimensionen des sprachlichen Raums. Variation – Mehrsprachigkeit – Konzeptualisierung*. Berlin etc., Lang: 39–59. (= *Schriften zur deutschen Sprache in Österreich* 45).
- Pickl, Simon et al. (2014): „Linguistic distances in dialectometric intensity estimation“. *Journal of Linguistic Geography* 2/1: 25–40.
- Pröll, Simon (2015): *Raumvariation zwischen Muster und Zufall. Geostatistische Analysen am Beispiel des Sprachatlas von Bayerisch-Schwaben*. Stuttgart: Steiner (= *Zeitschrift für Dialektologie und Linguistik, Beihefte* 160).
- Pröll, Simon/Elspaß, Stephan/Pickl, Simon (im Druck): „Areal microvariation in Germanspeaking urban areas (Ruhr Area, Berlin, and Vienna)“. In: Ziegler, Arne et al. (eds.): *Urban Matters. Current Approaches of International Sociolinguistic Research*. Amsterdam: Benjamins. (= *Studies in Language Variation*).
- Reichmann, Oskar (1988): „Zur Vertikalisierung des Varietätenspektrums in der jüngeren Sprachgeschichte des Deutschen“. Unter Mitwirkung von Christiane Burgi, Martin Kaufhold und Claudia Schäfer. In: Munske, Horst Haider et al. (eds.): *Deutscher Wortschatz. Lexikologische Studien. Ludwig Erich Schmitt zum 80. Geburtstag von seinen Marburger Schülern*. Berlin/New York, de Gruyter: 151–180.
- Rumpf, Jonas et al. (2009): Structural analysis of dialect maps using methods from spatial statistics. *Zeitschrift für Dialektologie und Linguistik* 76/3: 280–308.
- Séguy, Jean (1971): „La relation entre la distance spatiale et la distance lexicale“. *Revue de Linguistique Romane* 35: 335–357.
- Speelman, Dirk/Grondelaers, Stefan/Geeraerts, Dirk (2003): „Profile-based linguistic uniformity as a generic method for comparing language varieties“. *Computers and the Humanities* 37: 317–337. Statistik Austria (2020): statistik.at/web_de/klassifikationen/regionale_gliederungen/gemeinden/index.html [07.10.2021].
- Tagliamonte, Sali A./D’Arcy, Alexandra/Rodríguez Louro, Celeste (2016): „Outliers, impact, and rationalization in linguistic change“. *Language* 92/4: 824–849.
- Trudgill, Peter (1974): „Linguistic change and diffusion: description and explanation in sociolinguistic dialect geography“. *Language in Society* 3/2: 215–246.
- Trudgill, Peter (1986): *Dialects in Contact*. Oxford: Blackwell.

- van Gemert, Ilse (2002): *Het geografisch verklaren van dialectafstanden met een geografisch informatiesysteem (GIS)*. Masterarbeit, Rijksuniversiteit Groningen.
- Wardenga, Ute (2002): „Räume der Geographie und zu Raumbegriffen im Geographieunterricht“. *Wissenschaftliche Nachrichten* 120: 47–52.
- Wiesinger, Peter (1983): „Die Einteilung der deutschen Dialekte“. In: Besch, Werner et al. (eds.): *Dialektologie. Ein Handbuch zur deutschen und allgemeinen Dialektforschung*. Berlin/New York, de Gruyter: 807–900. (= *Handbücher zur Sprachund Kommunikationswissenschaft* 1.2).

Fußnote

- 1 Unter „Sprachraum“ verstehen wir – ähnlich wie Frings (1956: 24–26) – eine besondere Ausprägung des sozialen bzw. Kommunikationsraums, die nicht mit dem geographischen Raum identisch, aber durch ihn mitbedingt ist und durch sprachlich relevante Relationen unter einzelnen Orten geprägt sind, die sich als relationale Lagebeziehungen verstehen und so als räumliche Struktur beschreiben lassen (cf. Pickl 2013a: 56–63).
- 2 Diffusionen spielen sich in den Räumen ab, die die topologischen Modelle beschreiben, und schaffen und überformen dadurch ebendiese Räume, die wiederum die Basis (den „Container“) von Diffusionsentwicklungen darstellen. Dies ist kein konzeptioneller Mangel, sondern Ausdruck des fortdauernden Sprachwandels und der dadurch stattfindenden Überformung und Transformation des als elastisch zu verstehenden topologischen Sprachraums.
- 3 Verschiedene Studien stellen unterschiedlich starke Effekte der Berücksichtigung von Reisedistanzen gegenüber euklidischen Distanzen fest, was auch mit den jeweiligen topographischen Bedingungen zusammenzuhängen scheint: Van Gemert (2002) stellt bei der logarithmischen Regression mit linguistischen Distanzen im Verhältnis zu geographischer Distanz und Reisedistanz für niederländische Dialekte praktisch keine Veränderung fest, Nerbonne/Heeringa (2007) beobachten sogar eine leichte Verschlechterung; hier sind Reiseund geographische Distanzen mit $r = 0,92$ hoch korreliert. Dagegen identifiziert Gooskens (2004) für Norwegen eine deutliche Verbesserung von $r^2 = 0,17$ auf $0,29$, was unter anderem auf die extremere Topographie Norwegens zurückzuführen ist; hier beträgt die Korrelation zwischen geographischer und Reisedistanz lediglich $r = 0,68$. Szmrecsanyi (2012)
- 4 stellt mit Daten aus Großbritannien eine leichte Verbesserung der Vorhersagequalität bei Reisezeiten ($r^2 = 0,076$) gegenüber geographischen Distanzen fest ($r^2 = 0,044$).
- 5 Im Einzelfall sind solche Effekte jedoch durchaus zu beobachten (cf. Bailey et al. 1993; Britain 2004: 39), sodass hier von einer Anwendbarkeit des hierarchischen Modells nur mit negativem Vorzeichen ausgegangen werden kann. Mitunter wird für solche Fälle eigens das sogenannte kontrahierarchische Modell angesetzt („counter-hierarchical model“, Britain 2004: 39; 2012b: 455). Dieses Modell, das nicht als Gegenentwurf, sondern als konzeptionelle Ergänzung zum hierarchischen Modell angelegt ist, geht davon aus, dass bei einzelnen Innovationen die Richtung der Diffusion umgekehrt sein kann, i. e. dass Varianten aus ländlichen Gebieten in die Städte diffundieren und so die Hierarchie „hinaufklettern“.
- 6 Labov (2003: 9) nennt das hierarchische Modell „more general“, das Gravitationsmodell „more specific“ (cf. auch Labov 2007: 350).
- 7 Trudgills Gravitationsmodell wurde verschiedentlich Anwendungen unterzogen, um seine empirische Plausibilität zu überprüfen – mit

- unterschiedlichen Resultaten (cf. z. B. Chambers/Trudgill 1998: 178–185; Boberg 2000; Szmrecsanyi 2012).
- 8 Trudgill (1974: 233) verweist auf „Olsson 1965; and Haggett 1965: 35, for further discussion and treatment of – often serious – problems“.
 - 9 Die Quantifikation der Interaktion zwischen zwei Orten wird bei Trudgill (1974: 234–235) ergänzt um eine davon abgeleitete Modellierung des sprachlichen Einflusses von einem Ort auf einen anderen (unter der Annahme, „that interaction consists of influence in each direction proportional to population size“). Es handelt sich hier – anders als bei der Interaktion – um asymmetrische Beziehungen, da größere Orte stärkeren Einfluss auf kleinere haben als umgekehrt. Für die vorliegende Untersuchung ist jedoch ausschließlich die symmetrische Interaktion relevant, da asymmetrische Relationen nicht sinnvoll mit symmetrischen Relationen wie der linguistischen Distanz verglichen werden können (cf. auch Nerbonne/Heeringa 2007: 280; Szmrecsanyi 2012: 222).
 - 10 Hierbei wird auf die Einbindung des Proportionalitätsfaktors s des Gravitationsmodells verzichtet, da es sich dabei hier um ein Explanandum und nicht um ein Explanans handelt (cf. ähnlich Nerbonne/Heeringa 2007: 272; Szmrecsanyi 2012: 222).
 - 11 Wie auch beim Gravitationsmodell findet dies ohne Berücksichtigung möglicher sozialer Unterschiede, perzipierter Wertungen etc. statt. Diese dort mitunter kritisierte Vereinfachung stellt eine der Modellierung geschuldete Abstraktion dar.
 - 12 Ähnlich wie bei Trudgill (1974) wird auch bei Pickl (2013a) ein Maß für die sprachliche Ähnlichkeit in die Gleichung eingeführt, dort jedoch durch Ersetzen der geographischen Distanz d_{geo} durch die empirisch zu erhebende linguistische Distanz bzw. Unähnlichkeit d_{ling} . Aus demselben Grund, aus dem oben auf die Berücksichtigung von Trudgills Parameter s verzichtet wurde, wird auch hier von der Ersetzung von d_{geo} durch d_{ling} abgesehen.
 - 13 Für eine vergleichbare Untersuchung, in der die Vorhersagequalität von geographischer Distanz, Reisezeit, Interaktion (nach Trudgills Gravitationsmodell) sowie Dialektgebiete für die morphologische Distanz in englischen Dialekten überprüft wurde, cf. Szmrecsanyi (2012). Dort erzielten anhand einer Clusteranalyse derselben Daten ermittelte Dialektgebiete mit $r^2 = 0,325$ und einem logarithmischen Zusammenhang den höchsten Wert (wobei kritisch angemerkt werden muss, dass das Vorgehen hier etwas zirkulär war), geographische Distanzen den niedrigsten ($r^2 = 0,044$, linear; Reisedistanz: $0,076$, logarithmisch; Gravitationsmodell: $0,241$, logarithmisch).
 - 14 Zu näheren Angaben zur Methodik der Erhebung und zur Datenaufbereitung und -darstellung cf. Möller/Elspaß 2015.
 - 15 Dies ergibt bei 124 Ortspunkten $124 \cdot 124 = 15\,376$ paarweise Distanzen; die Distanzen zwischen Orten und sich selbst, die jeweils 0 betragen, nicht mitgerechnet, $15\,376 - 124 = 15\,252$ Distanzen.
 - 16 Es kommt auch andernorts durchaus vor, dass der lineare Zusammenhang stärker ist als der logarithmische (siehe z. B. Nerbonne/Heeringa 2007: 288f.; auch dort scheint visuell eher ein logarithmischer Zusammenhang nahezuliegen).
 - 17 Vermutlich trägt die oben beschriebene höhere alltagssprachliche Unähnlichkeit bei Distanzen über 300 km dazu bei, dass die lineare Regression hier besser abschneidet als die logarithmische.
 - 18 Nerbonne/Heeringa (2007) verwenden die multiplizierten Populationszahlen, i. e. das Quadrat des geometrischen Mittels und damit einen vergleichbaren Wert mit demselben Informationsgehalt.
 - 19 Obwohl die Korrelation mit logarithmierten Bevölkerungszahlen etwas niedriger ist als linear (siehe Tabelle 4), führt die logarithmische Darstellung in Abbildung 11 zu einer besser erkennbaren Werteverteilung.

- 20 Nerbonne/Heeringa (2007: 288) errechnen einen ..-Wert von 51,1 % für das Gravitationsmodell mit ihren niederländischen Daten. Allerdings scheinen sie nicht Trudgills Berechnungsmethode zu verwenden, sondern eine multiple Regression mit geographischen Distanzen und dem Produkt der Bevölkerungszahlen als unabhängigen
- 21 Variablen. Für einen vergleichbaren Ansatz siehe Abschnitt 4.4 in diesem Beitrag. Überraschenderweise stellen Nerbonne/Heeringa (2007: 287) eine positive Korrelation zwischen dem „population product“ und den linguistischen Distanzen fest – „exactly the opposite of what the gravity model predicts“.
- 22 Die jeweiligen konkreten Werte der einzelnen Parameter in den in Tabelle 8 vorgestellten Varianten des Regressionsmodells sind für die hier verfolgten Zwecke irrelevant und werden deshalb aus Gründen der Übersichtlichkeit nicht angegeben.
- 23 It is not suggested that all linguistic changes follow the same pattern of diffusion.“ (Labov 2003: 10).