

„Kann man das überhaupt messen?“ – Der Bedarf an detaillierten räumlichen Bevölkerungsdaten in der Wirtschaftsgeographie

Paula Prenzel

1 Einleitung

Belastbare Informationen zur Größe und Zusammensetzung der Bevölkerung sind Voraussetzung für eine Vielzahl sozialer, wirtschaftlicher und politischer Fragestellungen. Sowohl in der wissenschaftlichen Analyse von gesellschaftlichen Strukturen und Prozessen als auch in der Umsetzung von politischen Maßnahmen werden detaillierte sozio-ökonomische Daten benötigt. Damit diese als Grundlage für politische Entscheidungen oder empirische Forschungsergebnisse dienen können, müssen sie qualitativ hochwertig sein.

Dies ist besonders im Kontext der Sozialwissenschaften, die regelmäßig mit Problemen der Operationalisierung (also Messbarmachung) konfrontiert sind, von Bedeutung. Dort ist es der Normalfall, dass Phänomene nicht direkt beobachtet und höchstens imperfekt gemessen werden können. Ausnahmen, etwa gesellschaftliche Faktoren oder administrative Daten, die relativ verlässlich festzustellen sind, haben gerade deshalb einen hohen Stellenwert. Die demographischen Eckpunkte, die aus Bevölkerungsdaten abzuleiten sind, können so beispielsweise Grundlagen oder Kalibrierungsmöglichkeiten für Indikatoren abstrakterer Konzepte bieten.

Die Erhebung von verlässlichen Bevölkerungsdaten ist deshalb sehr wichtig, gleichzeitig aber methodologisch anspruchsvoll, und stellt, wie in diesem Sammelband zusammengefasst, auch große Herausforderungen für den Datenschutz dar. Dies trifft insbesondere zu, wenn die Daten granular genutzt werden sollen, also z.B. auf sub-nationaler geographischer Ebene oder mit weniger stark aggregierten Merkmalen. Somit ist die Verfügbarkeit von selbst sehr grundlegenden Bevölkerungsdaten in Deutschland mitunter überraschend eingeschränkt und die im Titel gestellte Frage, ob ein bestimmtes Konzept überhaupt messbar ist, von zentraler Bedeutung in der empirischen Forschung: Sie entscheidet, ob Fragestellungen umsetzbar sind und setzt unter Umständen der wissenschaftlichen Neugier und

dem Erkenntnisgewinn eine Grenze. Die empirische Forschung mit Bevölkerungsdaten befindet sich also in einem Spannungsfeld zwischen theoretischer und praktischer Relevanz auf der einen Seite und methodologischen sowie datenschutzrechtlichen Einschränkungen auf der anderen.

Ziel dieses Beitrages ist es, eine Nutzungsperspektive auf Bevölkerungsdaten, wie den Zensus, zu bieten. Aus Sicht der quantitativen (Wirtschafts-)Geographie wird diskutiert, warum detaillierte Bevölkerungsdaten für Planungs- und Forschungszwecke wichtig sind, welche Herausforderungen ihre Erhebung beinhaltet und welche Rolle der Zensus spielen kann. Danach werden einige Reflexionen zum Datenschutz in der Nutzung von Sekundärdaten beschrieben. Am Beispiel des Konzepts der *kulturellen Diversität* wird schließlich gezeigt, wie granulare räumliche Bevölkerungsdaten in der empirischen Forschung genutzt werden können. Dieser Beitrag erhebt dabei keinen Anspruch auf Vollständigkeit der möglichen Betrachtungsperspektiven, sondern will eine spezifische Nutzerperspektive abbilden, die sich grob im interdisziplinären Feld der Wirtschaftsgeographie einordnen lässt.

2 Die Nutzung von Bevölkerungsdaten

2.1 Anwendungsbereiche von Bevölkerungsdaten aus dem Zensus

Bevölkerungsdaten werden für ein breites Spektrum an Anwendungen benötigt, von denen hier am Beispiel des Zensus nur einige genannt werden sollen. Insgesamt stellt die Bevölkerungsgröße zweifellos eine zentrale Kennziffer gesellschaftlicher Belange dar, vor allem wenn es um Verteilungsfragen geht. Die Frage, wie viele Menschen eigentlich in Deutschland leben und wo, ist für politische, administrative, aber auch wirtschaftliche und soziale Zwecke wichtig. Bevölkerungszahlen werden in der Abgrenzung von Wahlkreisen berücksichtigt und spielen auch bei Gebietsreformen, also Änderungen der räumlichen Begrenzung von Gemeinden oder Kreisen, eine Rolle. Die Bevölkerungsgröße fließt außerdem in alle anderen Indikatoren ein, die relativ zur Einwohnerzahl gemessen werden, wie die Arbeitsproduktivität oder das Bruttoinlandsprodukt (BIP) pro Kopf, das zur Schätzung der relativen Wirtschaftskraft auf nationaler oder regionaler Ebene genutzt wird. Anhand des BIPs pro Kopf werden wirtschaftliche Disparitäten bewertet und zum Teil auch Förderfähigkeit für strukturpolitische Maßnahmen festgestellt (z.B. im Rahmen der EU-Kohäsionspolitik

oder, in Deutschland, der Gemeinschaftsaufgabe „Verbesserung der regionalen Wirtschaftsstruktur“ (GRW)). Auf sub-nationaler Ebene ist die Bevölkerungsgröße und -struktur außerdem essenziell für die Bedarfsanpassung im Rahmen der Raumplanung sowie für den öffentlichen Haushalt. Die regionale Infrastruktur- und Wohnraumplanung muss demographische Faktoren berücksichtigen, um sowohl die aktuelle als auch die langfristige Versorgung sicherzustellen. Gleichzeitig werden im Rahmen des Finanzausgleichs auf Ebene der Bundesländer und auch auf kommunaler Ebene Steuereinnahmen relativ zur Einwohnerzahl umverteilt, um Disparitäten in der Finanzkraft anzugleichen und eine flächendeckende Versorgung mit öffentlichen Diensten zu gewährleisten (vgl. den Beitrag von Fischer und Hufeld in diesem Sammelband).

Diese kurze Aufstellung zeigt bereits, dass Bevölkerungsdaten Kennzahlen mit weitreichenden Anwendungsgebieten liefern. Es wird dabei auch deutlich, dass die effektive Nutzung dieser Daten, zum Beispiel im Rahmen der Raumplanung oder des Finanzausgleichs, maßgeblich davon abhängt, dass die Bevölkerungsgröße tatsächlich richtig gemessen wird. In Deutschland wurde, bedingt durch den langen Abstand zu den letzten Volkszählungsergebnissen, die Bevölkerungszahl für über 20 Jahre anhand von Fortschreibungen geschätzt. Die Erfahrung des Zensus 2011 hat gezeigt, dass diese Methode zu einer Überschätzung der Bevölkerungsgröße führt, weil Abmeldungen in Meldeämtern nur unvollständig vorgenommen werden. Die Anpassung anhand des Zensus 2011 reduzierte die Bevölkerungszahl für Deutschland um fast 1,5 Millionen Einwohner*innen (1,8%) und führte vor allem auch auf Ebene der Bundesländer und Kommunen zu Verschiebungen (Ragnitz, 2013). Auch wenn die Auswirkungen dieser Korrektur für die öffentlichen Haushalte recht gering waren (Hennecke, 2013; Ragnitz, 2013), zeigt es dennoch die Notwendigkeit Bevölkerungsdaten regelmäßig und umfassend zu erheben.

Ein regelmäßiger Zensus führt nicht nur zur Korrektur des Bevölkerungsstandes und einer akkurateren Momentaufnahme der Bevölkerungsdaten, er kann außerdem genutzt werden, um die Zuverlässigkeit anderer Datenquellen zu verbessern. Der Zensus bietet die Grundlage für Stichproben, wie z.B. dem jährlichen Mikrozensus, in dem 1% der deutschen Bevölkerung befragt wird und der wiederum selbst zur Anpassung des Sozio-ökonomischen Panels (SOEP) genutzt wird (Wagner, 2010). So können auch zwischen den Volkszählungsjahren Daten mit größerer Präzision und Aussagekraft erhoben werden.

2.2 (Regionale) Bevölkerungsdaten in der Forschung

Bisher wurde der Nutzen von Bevölkerungsdaten eher im politischen und administrativen Sinne betrachtet. Ein weiterer zentraler Nutzungsbereich für diese Daten ist aber die empirische Forschung, die nun, vor allem aus der Perspektive der Wirtschaftsgeographie, thematisiert werden soll.

Wie bereits erwähnt, werden Bevölkerungsdaten genutzt, um sehr vielfältige Indikatoren zu erhalten, die sich nicht nur auf demographische Aspekte wie das Bevölkerungswachstum oder die Altersstruktur beschränken, sondern tatsächlich viele Maße umschließen, die Häufigkeiten in der Bevölkerung abbilden und deshalb relativ zur Einwohnerzahl berechnet werden. Diese fließen in wissenschaftliche Analysen in allen Sozialwissenschaften ein – ob als Kontrollvariable oder als Forschungsgegenstand. Außerdem sind vor allem auch andere Merkmale wissenschaftlich interessant, die gemeinsam mit der Einwohnerzahl erhoben werden können und differenziertere Aussagen über ihre Zusammensetzung zulassen. Dies betrifft weitere demographische Merkmale (z.B. Alter, Geschlecht, Staatsangehörigkeit), aber auch Informationen zum Bildungsstand oder Arbeitsmarktstatus. Je nach Datenquelle variieren die verfügbaren Merkmale, denn das Einwohnermeldeamt erhebt natürlich als Primärzweck andere Informationen als beispielsweise die Arbeitsämter. Außerdem können umfragebasierte Quellen viel detailliertere Informationen abbilden als administrative Daten, haben aber den Nachteil, dass sie meist nur für Stichproben umsetzbar sind. Aus wissenschaftlicher Perspektive sind deshalb einige im Zensus enthaltene Merkmale von besonderer Bedeutung, weil diese sonst nur schwierig in diesem Umfang erhoben werden können: Das betrifft beispielsweise Informationen zur Bildung, Nationalität und Religion, aber vor allem auch zum Wohnungsmarkt.

Obwohl Bevölkerungsdaten in der empirischen Forschung also eine zentrale Rolle spielen, setzt die wissenschaftliche Nutzung von Daten auch immer gewisse methodologische Anforderungen voraus. Daten können beispielsweise nur zuverlässig ausgewertet werden, wenn sie nicht durch unbekannte Faktoren verzerrt sind – eine Eigenschaft, die intuitiv oft als Repräsentativität beschrieben wird. Wenn die Daten für bestimmte Personen oder Gruppen zum Beispiel inkorrekt oder gar nicht erhoben wurden, kann dies zu systematischen Fehlern führen, die die statistische Auswertung verhindern oder verfälschen. Außerdem stellt die wissenschaftliche Analyse auch Ansprüche an die Zeitdimension von Daten. Auf der einen Seite werden für viele Zwecke Daten mit gewisser Aktualität benötigt, auf

der anderen Seite brauchen wir für die Untersuchung von Veränderungen aber auch Zeitreihen, also mehrere vergleichbare Erhebungen in regelmäßigem Abstand. Obwohl manche Fragestellungen anhand eines einzelnen Datensatzes bearbeitet werden können, müssen oft auch Daten aus verschiedenen Quellen zusammengetragen und richtig kombiniert werden. In der Geographie bedeutet dies häufig die Einbeziehung verschiedener räumlicher Ebenen. So werden beispielsweise Mikrodaten, also Daten für Individuen (z.B. zum Migrationshintergrund einer Person) genutzt, um Indikatoren auf regionaler Ebene zu berechnen (z.B. den Anteil der Bevölkerung mit Migrationshintergrund in Hamburg) oder Daten sub-nationaler Einheiten werden für verschiedene Länder kombiniert. Diese Operationen können aber nur gelingen, wenn Daten eindeutig im geographischen Raum zu verorten sind. Das bedeutet, dass die Zuordnung der Gebietseinheiten (z.B. im Rahmen der europäischen Einteilung der NUTS (*Nomenclature of Territorial Units for Statistics*)-Regionen transparent und soweit möglich konstant bleiben muss, damit auch die internationale und zeitliche Vergleichbarkeit gegeben bleibt.

Prinzipiell gibt es in der quantitativen empirischen Forschung zwei Möglichkeiten an Daten zu gelangen, die den methodologischen und inhaltlichen Ansprüchen einer Fragestellung gerecht werden: die Erhebung geeigneter Daten und die Nutzung von Sekundärdaten, d.h. Daten, die von anderen erhoben und eventuell bereits verarbeitet wurden. Leider ist ersteres, also die umfangreiche Selbsterhebung von Bevölkerungsdaten, für quantitative Studien auf sub-nationalen geographischen Ebenen meist unerreichbar: Die nötige Größenordnung der zu erhebenden Daten, Selektionsprobleme und die Erfassung von Zeitdimensionen stellen große Hürden dar. Schließlich müssten für eine Studie auf der Ebene der administrativen Kreise in Deutschland (wie in der Wirtschaftsgeographie üblich) Daten für 400 Regionen gesammelt (und jeweils die nötigen Stichprobengrößen und -verteilungen beachtet) werden. Doch auch Sekundärdaten sind nicht uneingeschränkt für geographische Analysen geeignet, denn die räumliche Ebene muss bereits in der Erhebung mitgedacht werden, damit ausreichende Stichprobengrößen in den sub-nationalen Gebietseinheiten zur Verfügung stehen. Selbst der Mikrozensus, der mit einer Stichprobengröße von 1% der deutschen Bevölkerung eine sehr große Haushaltsbefragung darstellt, ist unterhalb der geographischen Ebene der sogenannten Anpassungsschichten (etwa 123 räumliche Einheiten mit je ca. 500.000 Einwohnern), nur sehr eingeschränkt für Regionalanalysen nutzbar (Weil, 2009). Eine Aufbereitung des Mikrozensus zur Gewinnung von sehr klein-

räumigen Abbildungen (unterhalb der Gemeindeebene) ist aktuell nicht umsetzbar und selbst die Nutzung auf Kreis- oder Gemeindeebene ist methodologisch und organisatorisch schwierig, weil z.B. eine aufbereitete Regionalversion des Mikrozensus ausschließlich für das Jahr 2000 zur Verfügung steht (Pforr, 2021). Da die Nutzung von stichproben-basierten Befragungsdaten für räumlich tiefgegliederte Ebenen also sehr schwierig (bis unmöglich) ist, haben administrative Daten eine umso größere Bedeutung für quantitative Regionalstudien. Die wichtigsten Datenquellen sind so beispielsweise jene, die von den statistischen Ämtern und anderen Institutionen auf Kreis- oder gar Gemeindeebene zusammengestellt werden (und sich zum Teil auf den Zensus berufen), Arbeitsmarktbefragungen (z.B. im Rahmen der European Labour Force Survey) oder Daten aus den Sozialversicherungseinträgen, die vom Institut für Arbeitsmarkt- und Berufsforschung (IAB) der Bundesagentur für Arbeit aufbereitet werden. Alternativ werden natürlich auch andere Datenquellen genutzt, wie z.B. teilweise sehr detaillierte international verfügbare Daten, wie etwa Registerdaten in Skandinavien, die mit Gesundheitsdaten verknüpft sogar epidemiologische Kohortenstudien erlauben (Smith Jervelund & De Montgomery, 2020). Auch stellen, im Kontext von Big Data (vgl. den Beitrag von Gary Schaal in diesem Sammelband), kommerziell verfügbare Daten, Kooperationen mit Firmen sowie nicht zuletzt die innovative Nutzung online verfügbarer Informationen, wie Twitterdaten (z.B. Lansley & Longley, 2016), weitere nutzbare Datenquellen dar.

Die vorhergehende Diskussion hat gezeigt, dass Bevölkerungsdaten in der Wissenschaft sehr wichtig, aber nicht unbedingt leicht zu erhalten sind. Besonders wenn eine eigene Datenerhebung aus methodologischen oder finanziellen Gründen nicht umsetzbar ist, muss sich empirische Forschung in vielen Bereichen der Sozialwissenschaften mit gravierendem Datenmangel auseinandersetzen (Riphan, 2022). Eine Konsequenz unzureichender Datenverfügbarkeit ist, dass bestimmte Forschungsfragen überhaupt nicht oder nur in angepasster Form thematisiert werden können. So argumentiert *Haucap* (2022), dass sich die empirische Forschung aufgrund eingeschränkter Datenverfügbarkeit weniger auf Deutschland und mehr auf internationale Kontexte konzentriert, was wiederum zu einem Mangel an politisch verwertbaren wirtschaftswissenschaftlichen Erkenntnissen führen könnte.

Auch in der Geographie sind diese Probleme nicht zu unterschätzen und besonders bei fehlenden räumlichen oder nur stark aggregiert verfügbaren Daten spürbar. Wenn die Auswahl der zu untersuchenden regionalen

Gebietseinheiten aber stark durch die Datenverfügbarkeit gelenkt wird, ergeben sich im geographischen Kontext sowohl empirisch-methodologische als auch konzeptionelle Probleme. Methodologisch könnten z.B. Forschungsergebnisse maßgeblich von der Definition der räumlichen Einheiten beeinflusst werden, weil eine Studie vielleicht auf Kreisebene andere Ergebnisse zeigt, als auf Gemeindeebene, was in der Geographie als *Modifiable Areal Unit Problem* (MAUP) bezeichnet wird (Openshaw & Taylor, 1979; Madelin et al., 2009). Von dieser Problematik abgesehen, führt die Definition der geographischen Einheiten aufgrund von Datenverfügbarkeit außerdem zu einer konzeptionellen Kluft zwischen dem theoretischen und dem empirischen Raumverständnis: Obwohl ein geographischer Raum natürlich sozial konstruiert ist und sich konzeptionell nicht durch administrative Grenzen definieren lässt, ist z.B. das Raumkonzept relationaler Perspektiven der Wirtschaftsgeographie (z.B. Bathelt & Glückler, 2003) im Kontext eingeschränkter Datenverfügbarkeit nur in Ausnahmefällen quantitativ empirisch umzusetzen. Das bedeutet nicht unbedingt (wie mitunter vorgeworfen), dass eher regionalwissenschaftlich geprägte empirische Beiträge, die z.B. mit Daten auf Kreisebene arbeiten, konzeptionell eine starke räumliche Abgrenzung anhand administrativer Grenzen vertreten. Vielmehr ist die Nutzung von Daten dieser Raumeinheiten meist die einzige Möglichkeit bestimmte Prozesse überhaupt aus quantitativer geographischer Sicht zu untersuchen.

3 Zugang zu Forschungsdaten und Datenschutz

Datenschutzrechtliche Aspekte spielen in vielen wirtschaftswissenschaftlich-geprägten quantitativen Methoden nur eine untergeordnete Rolle, was vielleicht schon bezeichnend für die Problematik ist. Wie im Folgenden diskutiert wird, liegt das hauptsächlich in der Natur der typischen Methoden und der Nutzung von Sekundärdaten. Im Kontext von Datenerhebung (ob quantitativ oder qualitativ) und in anderen (Teil-)Disziplinen der Geographie und der Sozialwissenschaften insgesamt wird Datenschutz im Rahmen der Forschungsethik deutlich mehr betont. Außerdem sind Erklärungen zur Freiwilligkeit und informationellen Selbstbestimmung von Teilnehmenden in Befragungen, Interviews oder Experimenten, zum Zweck der Datenerhebung, -analyse und -speicherung, Zusicherung der Anonymität und Vertraulichkeit und andere Prinzipien des rechtlich und ethisch korrekten Umgangs mit Forschungsdaten natürlich, wie auch in den Leitli-

nien der DFG (2022) verankert, fundamental für die gute wissenschaftliche Praxis. Da es in diesem Beitrag aber um Bevölkerungsdaten im Kontext des Zensus geht, bei dem die Forschenden also die Daten nicht erheben, sondern nur nutzen, sollen hier einige Observationen zum Datenschutz in der Analyse von Sekundärdaten aus Nutzerperspektive skizziert werden.

Zunächst können zwei verschiedene Arten von Bevölkerungsdaten differenziert werden, für die die Datenschutzrelevanz unterschiedlich ausfällt: vergrößerte (aggregierte) Daten und Mikrodaten. Mit vergrößerten Daten sind hier Indikatoren gemeint, die zwar aus Individualbeobachtungen gewonnen wurden, die aber nur für ein beliebiges Aggregat (z.B. eine Gemeinde/Bundesland/Land) zur Verfügung gestellt werden. Je höher die Ebene der Vergrößerung und je häufiger das untersuchte Merkmal, desto weniger treffen Datenschutzprobleme zu. So würde wohl niemand behaupten, dass die Arbeitslosenquote auf Bundeslandebene datenschutzrechtlich bedenklich ist. Das liegt daran, dass man aus diesen stark aggregierten Daten nicht mehr auf Individuen schließen kann, selbst wenn der Kennzahl ursprünglich Informationen für einzelne Personen zu Grunde gelegen haben. Diese Art der Daten ist *absolut anonym* und es ist daher unbedenklich, wenn stark aggregierte Daten zu Forschungszwecken zur Verfügung gestellt werden. Aggregatdaten sind deshalb auch oft öffentlich zugänglich (wie z.B. die Datengrundlage für das Anwendungsbeispiel in Abschnitt 4).

Im Gegensatz zu Aggregatdaten liegen Mikrodaten auf Individualebene vor (z.B. einzelne Personen, Haushalte oder Firmen). Meist enthält jede einzelne Observation eines Mikrodatensatzes Informationen zu verschiedenen Merkmalen (wie z.B. Geschlecht, Alter und Bildungsstand). Da für eine bestimmte Person verschiedene Charakteristiken gleichzeitig betrachtet werden können, haben Mikrodaten ein deutlich größeres Nutzungspotenzial in wissenschaftlichen Untersuchungen, je detaillierter sie sind. Gleichzeitig sind sie natürlich aus Datenschutzperspektive auch bedenklicher, weil es um personenbezogene Daten geht, die prinzipiell tiefe Einblicke ermöglichen. Allerdings sind es eben auch genau diese Einblicke, die im Rahmen sozialwissenschaftlicher Fragestellungen untersucht werden sollen. Deshalb ist der Zugang und die Verarbeitung von Mikrodaten reglementiert und durch verschiedene Mechanismen gesichert (vgl. auch den Beitrag von Seckelmann in diesem Sammelband).

Das wichtigste Prinzip ist die Anonymisierung der Daten für Forschungszwecke. Bei Daten, die der *absoluten Anonymität* entsprechen, sind selbst Mikrodaten so umfassend vergrößert und kritische Merkmale entfernt, dass eine Identifikation der Personen nicht mehr möglich ist. Absolut

anonymisierte Daten können deshalb für verschiedene Zwecke als *Public-Use File* zur Verfügung gestellt werden. Aus wissenschaftlicher Perspektive sind absolut anonyme Daten aber meist nicht ausreichend.

Faktisch anonyme Mikrodaten dagegen sind Daten, bei denen prinzipiell eine De-anonymisierung nicht auszuschließen, aber doch sehr unwahrscheinlich ist. Diese Art der Daten, bei denen also eine Identifikation der Individuen höchstens durch sehr viel Aufwand möglich wäre, kann laut Bundesstatistikgesetz für wissenschaftliche Zwecke herausgegeben werden (nämlich als *Scientific-Use-File*). Allerdings geschieht dies nur unter strengen Auflagen. So kann die Nutzung solcher Mikrodaten nur zu Forschungszwecken an wissenschaftlichen Einrichtungen geschehen, die zum Teil (wie beispielsweise bei Mikrodaten von Eurostat) in einem separaten Prozess eine Akkreditierung erbitten müssen. Der Zugang zu den Daten muss beantragt werden und unterliegt letztendlich einem Datennutzungsvertrag mit der verwaltenden Organisation. Im Antrag müssen nicht nur das Forschungsvorhaben detailliert beschrieben, sondern auch alle Personen, die Zugang zu den Daten bekommen sollen, benannt, und in einem Datensicherheitskonzept der Ort und die technischen Voraussetzungen für die Speicherung und Analyse erklärt werden. Wird die Datennutzung genehmigt, wird ein Vertrag geschlossen, der die erlaubten Nutzungen der Daten (sowohl in der Speicherung als auch Verarbeitung und Publikation von Ergebnissen) sehr genau beschreibt und natürlich Versuche der De-anonymisierung sowie ungenehmigte Verknüpfung von Datensätzen untersagt. Die Datennutzung wird nur für einen beschränkten Zeitraum erlaubt, nach dem die Daten wieder vollständig gelöscht werden müssen. Bei manchen Datensätzen (vor allem jenen, in denen die faktische Anonymität eventuell nicht gegeben wäre) ist eine Nutzung nur in speziellen Forschungsdatenzentren möglich, d.h. Forschende müssen zur Analyse anreisen, dort unter strengen Bedingungen die Daten bearbeiten und dürfen nur datenschutzrechtlich unproblematische Ergebnisse wieder mitherausnehmen. In manchen Fällen ist auch eine Datenfernverarbeitung möglich, bei der Forschende die Syntax für ihre Analyse aus der Ferne übermitteln und die Ergebnisse dann zugesendet werden.

Selbst wenn Zugang zu einem *Scientific-Use-File* für relevante Mikrodaten besteht, bedeutet dies keinesfalls, dass Forschende nun alles über alle befragten Personen wissen oder überhaupt wissen wollen. Zum einen liegen in diesen faktisch anonymen Datensätzen einige Merkmale nur vergrößert vor, besonders im Kontext von Regionaldaten. In kleinen geographischen Gebieten können bestimmte Personengruppen nämlich so selten

vorkommen, dass sie dennoch zu identifizieren wären. Die faktische Anonymität wird also hier dadurch sichergestellt, dass bestimmte Informationen unterdrückt oder Gruppen nachträglich noch weiter zusammengefasst werden. So wird z.B. in den Daten der European Labour Force Survey die Staatsangehörigkeit nur anhand von großen geographischen Makroregionen (z.B. alle Mitgliedsstaaten der Europäischen Union) herausgegeben, sodass keine einzelnen Nationalitäten identifizierbar sind. Auf der anderen Seite geht es in quantitativen Analysen normalerweise darum, eine möglichst große und objektive Datenbasis anhand standardisierter Methoden zu untersuchen und dabei statistisch abgesicherte Aussagen zu Durchschnittswerten, Wahrscheinlichkeiten, Verteilungen und Muster im Datensatz zu identifizieren, mit denen theoretische Hypothesen überprüft werden können. Im Gegensatz zu Fallstudienansätzen, die vor allem auch in qualitativer Forschung üblich sind, beziehen sich quantitative Fragestellungen dabei meist nicht auf ein bestimmtes Individuum, denn statistische Methoden sind immer im Kontext von Unsicherheit und Wahrscheinlichkeit zu bewerten. Wir könnten zum Beispiel durch eine statistische Anwendung schließen, dass Körpergröße und Gewicht miteinander *korrelieren*, d.h. größere Menschen sind, durchschnittlich, schwerer als kleinere Menschen. Wir könnten sogar einen Faktor berechnen, der beschreibt wie viel mehr Menschen bei einem zusätzlichen Zentimeter Körpergröße durchschnittlich wiegen und könnten uns vielleicht, statistisch gesehen, relativ sicher in dieser Berechnung fühlen. Trotzdem würden wir beim Versuch, das Gewicht einer bestimmten Person nur anhand der Körpergröße vorherzusagen, in den meisten Fällen kläglich scheitern. Für den Kontext des Datenschutzes ist es deshalb relevant zu sehen, dass in der quantitativ-statistischen Forschung individuelle Daten zwar als Input dienen, Outputs, also Ergebnisse, sich aber normalerweise auf Durchschnitte beziehen. Quantitativforschende nutzen Individualdaten also im abstrakten Sinne als Teilkomponente einer großen Stichprobe und beschreiben in Ergebnissen keine Einzelfälle. Aber da Forschende oft eben genau an Merkmalen interessiert sind, die durch personenbezogene Daten repräsentiert werden, bleibt keine andere Wahl, als diese in den Quelldatensätzen so detailliert wie nötig (und rechtlich möglich) zu erhalten.

Die Hürden zum Datenzugang für selbst faktisch anonyme Daten sind sehr hoch, der Antrag und die Auflagen zeit- und zum Teil auch kostenintensiv. Manchmal ergeben sich aus Perspektive der Nutzenden auch überraschende Situationen, die direkte Konsequenzen für ein mögliches Forschungsdesign haben können. So stellt Deutschland im Rahmen der

europäischen Arbeitskräfteerhebung (European Labour Force Survey) zwar Informationen zur Staatsangehörigkeit der Befragten zur Verfügung, aber nicht (auch nicht vergrößert) zu den Geburtsländern, die allerdings im Zensus durchaus erhoben werden. Es ist schwer vorstellbar, dass dies technische Gründe hat, weshalb es vermutlich daran liegt, dass der rechtliche Rahmen eine Erhebung oder Übermittlung der Geburtsländer für die European Labour Force Survey nicht zulässt, obwohl zumindest auf den ersten Blick kein datenschutzrechtlicher Unterschied zu Informationen über die Staatsangehörigkeit zu erkennen ist und andere Länder beide Indikatoren zur Verfügung stellen. Empirisch macht es aber unter Umständen durchaus einen Unterschied, ob Informationen zum Geburtsland oder zur Staatsangehörigkeit betrachtet werden.

Insgesamt steht Forschung mit Bevölkerungsdaten also nicht nur vor der Schwierigkeit, dass es bestimmte Informationen einfach wirklich nicht gibt, Forschende sind auch gezwungen, Kosten-Nutzen-Abwägungen beim Zugang in ihre Forschungsdesigns einzubeziehen: Für manche Forschungsfragen lohnt sich der Aufwand, Daten zu beantragen oder gar eine mehrtägige Reise in ein Forschungsdatenzentrum zu unternehmen. Für andere Fragen (und je nach Mittelverfügbarkeit) könnte dies auch unüberwindbare Hindernisse darstellen, die zum Verwerfen wissenschaftlich und gesellschaftlich wertvoller Forschungsideen führen.

4 Das Beispiel wirtschaftlicher Effekte der Diversität

Diversität ist ein passendes Beispiel für ein Konzept, das sozialwissenschaftlich untersucht wird und sich fundamental auf Bevölkerungsdaten stützt. Es soll hier genutzt werden um eine konkrete Anwendung von Bevölkerungsdaten in der wirtschaftsgeographischen Literatur zu zeigen, die stark von der Verfügbarkeit granularer Daten abhängt. Natürlich wird Diversität oft im Kontext von Inklusion und Repräsentation genannt und hat wichtige soziologische, psychologische, rechtliche und politische Implikationen, auf die hier nicht weiter eingegangen wird. Stattdessen soll hier eine eher wirtschaftswissenschaftliche Perspektive eingenommen werden, die sich mit den Effekten und der Messung von Bevölkerungsdiversität, und insbesondere *kultureller Diversität*, beschäftigt.

Wie für viele sozialwissenschaftliche Konzepte stellt die Abgrenzung und Operationalisierung von Diversität eine große empirische Herausforderung dar. Denn auch wenn grundsätzlich eindeutig ist, dass Diversität eine Exis-

tenz von Vielfalt (zwischen Personen oder Gruppen) beschreibt, kann diese Vielfalt doch entlang verschiedener Dimensionen charakterisiert werden, die sich, im Sinne von Intersektionalität oder *Vertovecs* (2007) Begriff der „*Super-diversity*“ auch überschneiden. So könnten wir Diversität im Kontext einer langen Liste verschiedener Merkmale betrachten: ob z.B. Geschlecht, Alter, Nationalität, Bildung, Sprache, Religion, Einkommen oder sexuelle Orientierung – empirisch gesehen kann jede Art der Heterogenität im weitesten Sinne aus der Perspektive der Diversität untersucht werden.

In der wirtschaftswissenschaftlichen und wirtschaftsgeographischen Literatur wird der Begriff der Diversität oft im Migrationskontext betrachtet und dort auch als *kulturelle Diversität* bezeichnet. Konzeptionell greift die Operationalisierung von „Kultur“ anhand der nationalen Herkunft (z.B. durch die Staatsangehörigkeit oder das Geburtsland) natürlich viel zu kurz. Dennoch ist die Messung von Diversität anhand von Migrationscharakteristiken aus zwei Gründen wichtig. Erstens ist die nationale Herkunft vergleichsweise gut messbar und Daten sind (wie im Weiteren diskutiert) zwar nicht perfekt, aber für einige Anwendungen ausreichend. Zweitens hat Deutschland – wie auch andere Länder der Europäischen Union – schon seit langem steigende Immigration bei gleichzeitig fallenden Geburtsraten erfahren, was zu einer zunehmenden Internationalisierung der Bevölkerung führt (vgl. Coleman, 2008). Die damit einhergehenden Veränderungen, z.B. in der Erwerbsbevölkerung, haben auch wirtschaftliche Konsequenzen, die anfangs hauptsächlich im Kontext von Migrationseffekten, zunehmend aber auch in einem eigenen Forschungszweig zu ökonomischen Diversitätseffekten untersucht werden¹.

Es gibt verschiedene theoretische Ansätze, die einen wirtschaftlichen Effekt von Diversität vermuten lassen und hier nur sehr grob angeschnitten werden können. So stellt Immigration generell zum Beispiel einen Übertragungskanal für Humankapital, neues Wissen und Kompetenzen dar, was Produktivität und Innovation fördern kann (vgl. Breschi, Lissoni & Temgoua, 2016). Aber auch abgesehen von einem Zufluss neuer Kompetenzen durch Migration kann aus innovationsökonomischer Perspektive argumentiert werden, dass Vielfalt Innovation und damit Produktivitätsgewinne und Wachstum begünstigt. Kulturelle Vielfalt äußert sich beispielsweise auch durch eine Vielfalt an Wissen, Fähigkeiten, Erfahrungen, Lösungsansätzen, Netzwerken oder Ideen. Wenn wir Innovation, nach *Schumpeter*

1 Die hier im Folgenden benannten Texte sind nur exemplarisch zu sehen. Für detailliert Übersichtsartikel siehe Niebuhr und Peters (2020) oder Ozgen (2021).

(1934), als Prozess der Rekombination von existierenden Faktoren verstehen, vergrößert Diversität also die Bandbreite der Möglichkeiten: je vielfältiger das Wissen, desto mehr neue Rekombinationen sind möglich. Gleichzeitig wissen wir unter anderem aus der Literatur zu Diversität in Teams (z.B. Stahl et al., 2010), dass Vielfalt auch eine Herausforderung darstellen kann, z.B. wenn Kommunikationsschwierigkeiten oder Konfliktpotential auftreten, und damit nicht automatisch zu wirtschaftlichen Gewinnen führen muss.

Der wirtschaftliche Effekt von Diversität bleibt also eine empirische Frage, die aktuell in vielen unterschiedlichen Kontexten untersucht wird. So gibt es beispielsweise Belege für Korrelationen zwischen Diversität und wirtschaftlicher Leistung bzw. Wachstum (z.B. Alesina, Harnoss & Rapoport; Docquier et al., 2020), Produktivität (z.B. Trax et al., 2015; Kemeny & Cooke, 2018), zunehmendem Unternehmertum (Audretsch et al., 2010; Rodríguez-Pose & Hardy, 2015) und Innovation (z.B. Niebuhr, 2010, Brixey et al., 2020). Allerdings gibt es in der Literatur auch Hinweise für fehlende oder gar negative Diversitätseffekte, weshalb *Ozgen* (2021) die empirische Evidenz insgesamt als „gemischt“ beschreibt und den Bedarf an weiteren Studien betont, die insbesondere die kausalen Mechanismen weiter untersuchen. Eine zentrale Herausforderung für diese empirischen Ansätze bleibt aber die Verfügbarkeit der nötigen detaillierten Bevölkerungsdaten.

Anhand des Beispiels der Diversitätsdaten lassen sich einige der oben beschriebenen Punkte zum Datenschutz und der Datenverfügbarkeit noch einmal konkret wiederholen und weiter ausführen. Erstens sind die benötigten Daten für Diversitätsmaße personenbezogene und auch sensible Merkmale, bei denen der Datenschutz natürlich einen großen Stellenwert einnimmt. Gleichzeitig ist, zweitens, Diversität eine Eigenschaft der Bevölkerung, nicht des Individuums. Die relevante Bevölkerung muss entweder durch Organisationen (Firmen, Teams) oder geographisch abgegrenzt werden, wobei wichtig ist, dass genügend Observationen bleiben, um eine statistische Analyse zu ermöglichen. Drittens muss Diversität durch die Auswahl und Kombination der richtigen Indikatoren operationalisiert werden. So ist ein sehr grundlegender Indikator für die kulturelle Diversität der Anteil der Bevölkerung, der im Ausland geboren wurde oder nicht die Staatsbürgerschaft des Wohnlandes hält. Besonders der Anteil der im Ausland geborenen Bevölkerung wird häufig als einfacher Diversitätsindikator genutzt, da er eine eigene Migrationserfahrung und so einen Hinweis auf im Ausland erworbenes Wissen oder Erfahrungen bietet. Allerdings kann man

argumentieren, dass der Anteil der im Ausland geborenen Personen die Diversität nur unzureichend misst. Wenn in einer Region viele Menschen aus genau einem anderen Land wohnen, würden wir die Bevölkerung (die dann „nur“ durch zwei nationale Herkünfte charakterisiert ist) nicht für sehr divers halten. Deshalb betrachten komplexere Diversitätsindikatoren nicht nur die Anzahl der im Ausland Geborenen, sondern auch die relative Verteilung über verschiedene Geburtsländer. *Fraktionalisierungsmaße*², die in der Literatur oft angewendet werden, untersuchen, wie gleichmäßig die Bevölkerung über verschiedene Gruppen (z.B. Geburtsländer) verteilt ist und bilden damit intuitiv die Wahrscheinlichkeit ab, dass zwei zufällig ausgewählte Personen in der Bevölkerung im selben Land geboren wurden. Je mehr verschiedene Geburtsländer vorkommen und je gleichmäßiger die relativen Anteile über die Bevölkerung verteilt sind, desto größer die kulturelle Diversität. Die Berechnung von Fraktionalisierung, also einem präziseren und detaillierteren Indikator für Diversität, benötigt aber granulare Informationen zur Verteilung der Bevölkerung über verschiedene Gruppen (also Geburtsländer). Da aber, besonders in Regionen mit wenigen Personen aus bestimmten Ländern, die Information zu den Geburtsländern oft nicht einmal faktisch anonymisiert werden kann, ist die Verfügbarkeit stark eingeschränkt und z.B. in öffentlichen administrativen Daten schon für die Kreisebene (NUTS-3) nicht gegeben. Studien der regionalen Diversität für Deutschland werden deshalb entweder auf stichproben-basierten Befragungen oder anhand Arbeitsmarktdaten (z.B. der Integrierten Erwerbsbiografien des IAB) berechnet (wie beispielsweise in Niebuhr (2010) oder Brixy et al. (2020)), zu denen der Zugang aber, wie oben beschrieben, stark reglementiert ist.

Um diese Punkte nun noch graphisch zu präsentieren, zeigt Abbildung 1, welche Möglichkeiten der Diversitätsmessung es auf regionaler Ebene mit den öffentlich verfügbaren Daten des Zensus 2011 gibt. Die Datenbasis bezieht sich auf die Anzahl der Personen nach Geburtsland, die aus der Zensusdatenbank heruntergeladen und aufbereitet wurden. Die kleinräumigste Aufteilung, die dort angeboten wird, entspricht in etwa der NUTS-2 Ebene (die aus (ehemaligen) Regierungsbezirken und einigen Bundesländern besteht). Nur für Rheinland-Pfalz werden die Daten nicht anhand der korrekten NUTS-2-Aufteilung, sondern stattdessen als Bundesland zur Verfügung gestellt. Im oberen Teil von Abbildung 1 wird nur der grobe

2 Die Fraktionalisierung ist für Region i und Anteil der Bevölkerung (s) aus n Geburtsländern definiert als: $F_i = \sum_{k=1}^n s_{ik}^2$.

Indikator der im Ausland geborenen Bevölkerung (in Prozent) betrachtet, allerdings für verschiedene geographische Ebenen der Aggregation. Im nationalen Durchschnitt (Abb.1, oben links) waren laut Zensus 2011 13,2% der deutschen Bevölkerung außerhalb Deutschlands geboren, auch wenn, wenig überraschend, nicht räumlich gleich verteilt. Die Stadtstaaten, Nordrhein-Westfalen, Hessen und Baden-Württemberg liegen deutlich über dem nationalen Durchschnitt, wohingegen die ostdeutschen Flächenländer und Schleswig-Holstein darunter liegen (Abb.1, oben Mitte). Aber auch die Bundeslandebene versteckt starke räumliche Unterschiede: wenn wir die Daten noch etwas detaillierter betrachten (Abb. 1, oben rechts), wird klar, dass zum Beispiel nicht ganz Bayern als überdurchschnittlich divers eingeschätzt werden kann. Die im Ausland geborene Bevölkerung konzentriert sich in Bayern auf Oberbayern (mit 17,2%) und Mittelfranken (16,7%), die Bezirke mit den zwei größten Städten (München, Nürnberg). Im nationalen Vergleich ist der Norden und Osten Bayerns (gemessen am Anteil der im Ausland geborenen Bevölkerung) nicht besonders divers. Kleinere geographische Gebiete sind anhand dieser (absolut anonymen und öffentlich verfügbaren) Datenquelle nicht möglich und präzisere Muster (wie z.B. mögliche Stadt-Land-Unterschiede) sind nicht abzubilden.

Die unteren beiden Karten in Abbildung 1 zeigen die Diversität anhand der Fraktionalisierung der im Ausland geborenen Bevölkerung. Personen, die in Deutschland geboren sind, sind in diesem Indikator ausgenommen, da sie die (mit Abstand) größte Gruppe darstellen und den Indikator sonst stark beeinflussen würden. Die erste Karte (Abb. 1, unten links) zeigt, auf Basis der kleinstmöglichen geographischen Einheiten, die Diversität, wenn die unterschiedlichen Herkunftsländer in 12 große Makroregionen unterteilt werden, wie es in den vergrößerten Daten der European Labour Force Survey der Fall ist³. Die zweite Karte (Abb. 1, unten rechts) wertet dagegen den vollen Umfang der Geburtslanddaten aus, was hier 203 Gruppen entspricht⁴. Die Färbung der Karten unten ist anhand von vier Quantilen

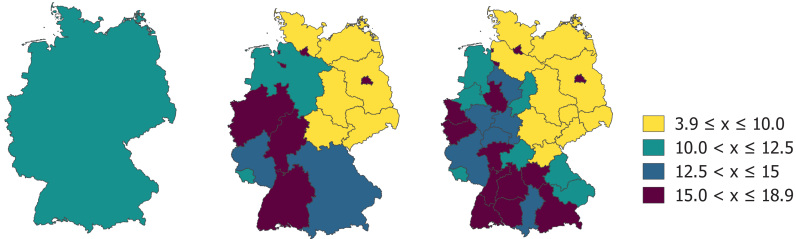
3 Die hier implementierten Makroregionen orientieren sich grob an den Gruppen, die in den 2004-2019 Versionen der European Labour Force Survey genutzt werden. Sie entsprechen: EU, EFTA, Other Europe, North Africa, Other Africa, Near and Middle East, East Asia, South and South East Asia, North America, Central America (and Caribbean), South America, Australia and Oceania.

4 Die Daten enthalten auch historische Geburtsländer, wie z.B. die Sowjetunion als Herkunftsland. Diese wurden so einbezogen. Observationen ohne Angabe sowie Staatenlose wurden für bessere Vergleichbarkeit ausgeschlossen.

vorgenommen, sodass jeweils die 25% der diversesten Regionen in beiden Karten die dunkelste Farbe hat.

Der Vergleich der Karten zeigt, erstens, dass die Vergröberung der Daten deutliche Auswirkungen auf die relative Verteilung der gemessenen Diversität hat. Wenn die Herkunftsländer als Makroregionen genutzt werden, scheinen die niedersächsischen Regierungsbezirke zum Beispiel diverser als die bayerischen, doch dieser Eindruck ist umgekehrt, wenn die Länder separat in die Fraktionalisierung einfließen. Zweitens fällt auf, dass es deutliche Unterschiede zwischen dem groben Anteil an im Ausland geborener Bevölkerung und der Fraktionalisierung gibt. So leben in Leipzig zwar insgesamt wenig im Ausland geborene Menschen, diese kommen aber aus sehr vielen verschiedenen Ländern, was wiederum auf eine recht hohe Diversität hinweisen kann. Natürlich ist nicht klar, welche der Karten nun amtreffendsten das abstrakte Konzept der kulturellen Diversität abbildet. Jedoch führt jeder Schritt der Vergröberung (ob auf geographischer Ebene oder für die Herkunftsländer) zu einem Verlust an Details und deshalb möglicherweise zu problematischen Verzerrungen oder gar unterschiedlichen empirischen Ergebnissen. Dies ist besonders kritisch, wenn die Daten nur vergrößert zur Verfügung stehen, und die Forschung nicht einmal die Sensitivität von Ergebnissen für verschiedene Einschränkungen überprüfen kann.

Im Ausland geborene Bevölkerung (in %)



Diversität der im Ausland geborenen Bevölkerung: Gruppen (links), Länder (rechts)

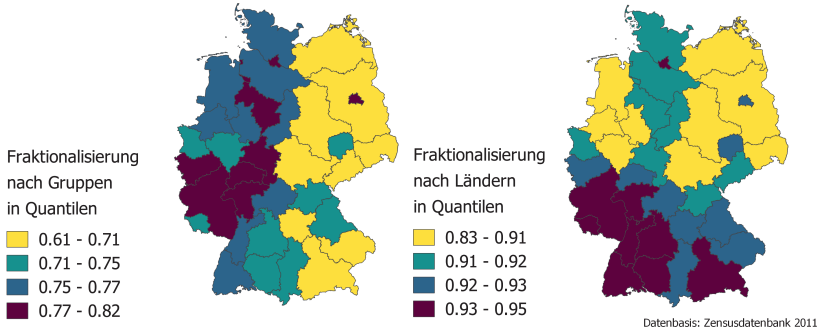


Abbildung 1: Darstellungen zur Diversität auf nationaler, Bundesland- und NUTS2-Ebenen. Eigene Darstellung auf Datenbasis der Zensusdatenbank (Statistische Ämter des Bundes und der Länder, 2022).

5 Fazit

Ziel dieses Beitrags war es, aus wirtschaftsgeographischer Perspektive zu diskutieren, welche Anwendungen und Herausforderungen die Nutzung von Bevölkerungsdaten, wie jene, die im Zensus erhoben werden, beinhaltet. Informationen zur Bevölkerung sind nicht nur administrativ und politisch zentral, sie sind auch aus den Sozialwissenschaften nicht wegzudenken. Dabei stellt die wissenschaftliche Verwertung aber bestimmte Ansprüche an Daten, weshalb die Erhebung von selbst grundlegenden Merkmalen zur Bevölkerungszusammensetzung methodologisch nicht unproblematisch ist. Besonders, wenn Daten flächendeckend auf sub-nationaler geographischer Ebene für quantitative Fragestellungen genutzt werden

sollen, ist es für Forschende meist nicht möglich, diese selbst zu erheben. Stattdessen muss sich die Forschung oft auf Daten verlassen, die auch für administrative Zwecke genutzt und deshalb in großem Stil gesammelt und verwaltet werden. Der Zensus, als regelmäßige und umfassende Erfassung von Bevölkerungsdaten, ist ein wichtiges Instrument zur Erhebung und Qualitätssicherung von Daten, die fundamental beschreiben, wer eigentlich wo in Deutschland lebt. Das heißt natürlich nicht, dass der Zensus tatsächlich die Realität perfekt abbildet, denn auch in Volkszählungen gibt es Ungenauigkeiten und Fehler (Wagner, 2010). Dennoch ist der Zensus, gerade im Vergleich zur Fortschreibung des Bevölkerungsstandes, zumindest geeignet eine Vergleichsgröße darzustellen und wegen der Auskunftspflicht auch Informationen zu erheben, die so umfassend sonst überhaupt nicht verfügbar sind (z.B. zur Wohnsituation).

Bevölkerungsdaten stellen einen Datenschatz dar, der natürlich für wissenschaftliche, aber eben potentiell auch für schädliche Zwecke genutzt werden kann. Deshalb wird ihre Nutzung im Rahmen des Datenschutzes streng überwacht. Dies bedeutet nicht nur, dass der Zugang zu Forschungsdaten reglementiert ist, sondern auch, dass die faktische Anonymität normaler Scientific-Use-Files eine starke Vergrößerung von bestimmten Merkmalen voraussetzt. Auch wenn dies aus Perspektive des Datenschutzes richtig ist, müssen wir uns bewusst sein, dass diese Einschränkungen Konsequenzen für die Forschung und damit den Erkenntnisgewinn haben. Denn es ist in der sozialwissenschaftlichen Forschung, wie auch in den Naturwissenschaften, wichtig, dass Informationen möglichst detailliert zur Verfügung stehen, weil mit jeder Vergrößerung der Daten (ob im geographischen oder inhaltlichen Sinne) Informationen verloren gehen. Für manche Anwendungen kann dieser Informationsverlust verkraftbar sein, für andere führt er im schlimmsten Fall zu einer Einschränkung der Genauigkeit oder Belastbarkeit der Ergebnisse. So ist zum Beispiel aus wirtschaftsgeographischer Perspektive die Verfügbarkeit granularer räumlicher Daten, die benötigt werden, um regionale Unterschiede in Deutschland zu untersuchen, nach wie vor sehr schlecht.

Die Verfügbarkeit von Daten ist aus empirischer Sicht nicht optional: Nur aus einer ausreichenden Datenbasis kann ein empirischer Erkenntnisgewinn entstehen. „Kann man das eigentlich messen?“ ist deshalb eine bezeichnende Frage in den Sozialwissenschaften. Man kann sie mit Bezug auf die Schwierigkeit, sozialwissenschaftliche Konzepte überhaupt zu operationalisieren, verstehen. Aber selbst wenn wir prinzipiell wissen, wie wir zum Beispiel Diversität messen wollen, bleibt dennoch die Frage der Messbar-

keit im Sinne der Datenverfügbarkeit. Sozialwissenschaftliche Fragestellungen haben gesellschaftliche Relevanz, denn sie erlauben uns zum Beispiel tatsächliche Prozesse zu beschreiben, Praktiken zu evaluieren und Politik auszurichten. Gerade im Vergleich mit den Möglichkeiten der Datenauswertung, die viele private Akteure haben (denen wir alle freiwillig und unkritisch selbst sensibelste Dinge zur Verfügung stellen), steht die wissenschaftliche Nutzung von öffentlichen Daten oft vor großen Hürden. Wenn Datenmangel oder Schwierigkeiten des Zugangs dazu führen, dass sich die Wissenschaft mit einigen Fragen nicht auseinandersetzen kann oder dass sie sich international andere Anwendungsbeispiele und Datenquellen sucht, schränken wir so auch die wissenschaftlichen Entdeckungen ein. Es ist deshalb wichtig, die wissenschaftliche Nutzung von Bevölkerungsdaten als legitimen und wertvollen Zweck anzuerkennen und neue Wege zu finden, wie sich diese Nutzung mit den hohen Datenschutzansprüchen in Deutschland vereinbaren lässt.

6 Literatur

- Alesina, A., Harnoss, J., & Rapoport, H. (2016). Birthplace diversity and economic prosperity. *Journal of Economic Growth*, 21, 101-138. <https://doi.org/10.1007/s10887-016-9127-6>
- Audretsch, D., Dohse, D., & Niebuhr, A. (2010). Cultural diversity and entrepreneurship: A regional analysis for Germany. *Annals of Regional Science*, 45(1), 55-85. <https://doi.org/10.1007/s00168-009-0291-x>
- Bathelt, H. & Glückler, J. (2003). Toward a relational economic geography. *Journal of Economic Geography*, 3, 117-144. <https://academic.oup.com/joeg/article/3/2/117/997236>
- Breschi, S., Lissoni, F., & Temgoua, C. N. (2016). Migration and Innovation: A Survey of Recent Studies. In R. Shearmur, C. Carrincazeaux, & D. Doloreux (Eds.), *Handbook on the Geographies of Innovation* (pp. 382-398). Cheltenham, UK: Edward Elgar.
- Brixy, U., Brunow, S., & D'Ambrosio, A. (2020). The unlikely encounter: Is ethnic diversity in start-ups associated with innovation? *Research Policy*, 49(4), 103950. <https://doi.org/10.1016/j.respol.2020.103950>
- Coleman, D. (2008). The demographic effects of international migration in Europe. *Oxford Review of Economic Policy*, 24(3), 453-477. <https://doi.org/10.1093/oxrep/grn027>
- Deutsche Forschungsgemeinschaft. (2022). *Leitlinien zur Sicherung guter wissenschaftlicher Praxis*. Kodex. Bonn: DFG. <https://doi.org/10.5281/zenodo.6472827>

- Docquier, F., Turati, R., Valette, J., & Vasilakis, C. (2020). Birthplace diversity and economic growth: evidence from the US states in the Post-World War II period. *Journal of Economic Geography*, 20, 321–354. <https://doi.org/10.1093/jeg/lbz016>
- Fischer, L. & Hufeld, U. (2024). Zensus und Steuern – Fiskalzensus: Informationsgewinnung und Datenverarbeitung im Finanz- und Steuerrecht. In D.-S. Valentiner & K. Goldberg (Hrsg.) *Zensus 2022 – Volkszählungen zwischen Recht, Politik und Sozialwissenschaft*, (S. 147-174). Baden-Baden: Nomos.
- Haucap, J. (2022). Datenmangel und andere Probleme der wirtschaftswissenschaftlichen Politikberatung in Deutschland. *Wirtschaftsdienst*, 102(7), 506-510. <https://doi.org/10.1007/s10273-022-3242-0>
- Hennecke, H. (2013). Auswirkungen der Zensusergebnisse auf den kommunalen Finanzausgleich. *Der Landkreis*, 83(6), 225-226.
- Kemeny, T., & Cooke, A. (2018). Spillovers from immigrant diversity in cities. *Journal of Economic Geography*, 18(1), 213–245. <https://doi.org/10.1093/jeg/lbx012>
- Lansley, G., & Longley, P. A. (2016). The geography of Twitter topics in London. *Computers, Environment and Urban Systems*, 58, 85-96. <https://www.sciencedirect.com/science/article/pii/S0198971516300394?via%3Dihub>
- Madelin, M., Grasland, C., Mathian, H., Sanders, L. & Vincent, J. (2009). Das “MAUP”: Modifiable Areal Unit – Problem oder Fortschritt? *Informationen zur Raumentwicklung*, 10, 645-660.
- Niebuhr, A. (2010). Migration and innovation: Does cultural diversity matter for regional R&D activity? *Papers in Regional Science*, 89(3), 563–585. <https://doi.org/10.1111/j.1435-5957.2009.00271.x>
- Niebuhr, A., & Peters, J. C. (2020). Population Diversity and Regional Economic Growth. In M. M. Fischer & P. Nijkamp (Eds.), *Handbook of Regional Science* (pp. 1–17). Berlin: Springer.
- Openshaw, S., & Taylor, P.J. (1979). A million or so correlation coefficients: Three experiments on the modifiable areal unit problem. In N. Wrigley (ed.) *Statistical Applications in the Spatial Sciences*, 127–144. London: Pion.
- Ozgen, C. (2021). The economics of diversity: Innovation, productivity and the labour market. *Journal of Economic Surveys*, 35(4), 1168–1216. <https://doi.org/10.1111/joes.12433>
- Pforr, K. (2021). *Regionale Kontextdaten mit dem Mikrozensus*. (GESIS Papers, 2021/02). Köln: GESIS - Leibniz-Institut für Sozialwissenschaften. <https://doi.org/10.21241/ssoar.71319>
- Ragnitz, J. (2013). Zensus 2011 und Finanzausgleich: Kein Grund zur Aufregung, *Wirtschaftsdienst*, 93 (7), pp. 426-427. <https://doi.org/10.1007/s10273-013-1545-x>
- Riphahn, R.T. (2022). „Wir wissen in Deutschland vieles nicht, was wir wissen sollten“: Ein Gespräch über die unzureichende Verfügbarkeit von Daten für die Forschung, die Minijob-Falle und die Vererbbarkeit von Sozialhilfeabhängigkeit in Deutschland. *Perspektiven der Wirtschaftspolitik*, 23 (1) pp. 38-48. <https://doi.org/10.1515/pwp-2022-0008>

- Rodríguez-Pose, A., & Hardy, D. (2015). Cultural diversity and entrepreneurship in England and Wales. *Environment and Planning A*, 47(2), 392–411. <https://doi.org/10.1068/a130146p>
- Schaal, G. S. (2024). Big Data ergänzt Zensus-Daten: Ein Paradigma für republikanisch inspirierte Governance in der digitalen Konstellation. In D.-S. Valentiner & K. Goldberg (Hrsg.) *Zensus 2022 – Volkszählungen zwischen Recht, Politik und Sozialwissenschaft*, (S. 89-108). Baden-Baden: Nomos.
- Schumpeter, J.A. (1934). *The theory of economic development*. Cambridge, MA: Harvard University Press.
- Seckelmann, M. (2024). Wissen und Recht: Die Nutzung der Zensus-Daten in der Rechtswissenschaft. In D.-S. Valentiner & K. Goldberg (Hrsg.) *Zensus 2022 – Volkszählungen zwischen Recht, Politik und Sozialwissenschaft*, (S. 133-146). Baden-Baden: Nomos.
- Smith Jervelund, S., & De Montgomery, C. J. (2020). Nordic registry data: value, validity and future. *Scandinavian Journal of Public Health*, 48(1), 1-4. <https://doi.org/10.1177/1403494819898573>
- Stahl, G. K., Maznevski, M. L., Voigt, A., & Jonsen, K. (2010). Unraveling the effects of cultural diversity in teams: A meta-analysis of research on multicultural work groups. *Journal of International Business Studies*, 41(4), 690–709. <https://doi.org/10.1057/jibs.2009.85>
- Statistische Ämter des Bundes und der Länder (2022). *Ergebnisse des Zensus 2011: Bevölkerung kompakt - Personen: Geburtsländer (1000A-1012)*. Wiesbaden: Statistisches Bundesamt. <https://ergebnisse2011.zensus2022.de/datenbank/online>
- Trax, M., Brunow, S., & Suedekum, J. (2015). Cultural diversity and plant-level productivity. *Regional Science and Urban Economics*, 53, 85–96. <https://doi.org/10.1016/j.regsciurbeco.2015.05.004>
- Vertovec, S. (2007). Super-diversity and its implications. *Ethnic and Racial Studies*, 30(6), 1024-1054. <https://doi.org/10.1080/01419870701599465>
- Wagner, Gert G. (2010). Zensus 2010/11: eine längst überfällige Erhebung, *DIW Wochenbericht*, 77(4), pp. 11-14.
- Weil, S. (2009). Regionaldaten der amtlichen Statistik – Ein Angebot im Spannungsfeld rechtlicher Restriktionen, statistisch-methodischer Beschränkungen und dem Bedarf an regional tief gegliederten Daten. *Statistische Monatshefte Rheinland-Pfalz*, 2009 (04). Bad Ems: Statistisches Landesamt Rheinland-Pfalz.

