

Heft 3

Epidemiologische Grundlagen der Sozialmedizin

Autoren:

Jens-Uwe Niehoff, Wolfgang Hoffmann

Vorwort zu Heft 3

Die Epidemiologie ist sowohl als eigenständige Wissenschaft wie auch mit ihrem Methodenarsenal ein unverzichtbarer Teil der medizinischen Forschungslandschaft. Sie ist auch für die Sozialmedizin Teil ihrer essenziellen Grundlagen.

Etwa seit den 1950er Jahren hat die Nutzung epidemiologischer Methoden eine erhebliche Erweiterung erfahren. Diese bezieht sich zunächst darauf, dass sich das Nutzungsfeld auf alle Krankheiten erweiterte, also nicht mehr nur für die quantitative Analytik übertragbarer Krankheiten und für die Vektoranalytik genutzt wurde. Sie fand nun in großer Breite auch in der Expositionsforschung vor allem in der Sozial-, Arbeits- und Umweltmedizin Anwendung.

Die Epidemiologie wurde gleichsam zur Schlüsselwissenschaft bei der Identifikation von Gefährdungsunterschieden zwischen Gruppen von Personen in unterschiedlichen Lebenssituationen, mit unterschiedlichen Lebensweisen und Gesundheitsressourcen. „Mögliches“ zu identifizieren, es zu fördern oder es zu verhindern wurde so zu einem eigenständigen Handlungsfeld. Risiko und Chance wurden Gegenstände systematischer Forschung und öffneten das Feld der Prädiktion.

Allerdings gab es neben der Medizin weitere Interessenten, hier vor allem die nach Prinzipien der Risikoäquivalenz agierende Kranken- und Lebensversicherungswirtschaft. Epidemiologie wurde so auch ein Konfliktfeld der Akteure im pro und contra von fiskalischer, parafiskalischer (hier vor allem solidarischer) sowie risikoselektiver und risikoäquivalenter Krankenversicherungen.

Eine erhebliche Erweiterung erfuhr das Arbeitsfeld der Epidemiologie durch die großen bevölkerungsbasierten Interventionsstudien der 1970er Jahren in den USA. Sie folgten der Erwartung, es ließen sich mit prädiktiven Interventionen alle jene Sterbeursachen bekämpfen, die bei gestiegener mittlerer Lebensdauer und im Tausch gegen die „classic killers“, vor allem die Tuberkulose und die Ursachen der Säuglings- und Kindersterblichkeit, nun die Todesursachenlisten anführten. Unvermeidbar wurde die Epidemiologie so auch zu einem Brennpunkt in den Auseinandersetzungen um gesundheits- und versorgungspolitische Interessen zwischen den Polen Mengenausweitung und Begrenzung auf diagnostische und therapeutische Interventionen nach den Maßstäben von Notwendigkeit, Angemessenheit und Wirtschaftlichkeit.

Vor diesem Hintergrund erweiterte sich das Verständnis der Epidemiologie auf alle quantitativen Studien in der Medizin, vor allem auch auf klinische Studien. In der Folge verbreitete sich eine Sicht, die jede Anwendung mathematisch-statistischer Methoden dann auch der Epidemiologie zurechnete.

Diese Sicht wird hier nicht geteilt. Konstitutiv für das Verständnis des Forschungsfeldes „Epidemiologie“ bleibt hier die Sicht, dass es sich um die Lehre von der Verteilung von Krankheiten in realen Bevölkerungen, den Ursachen dieser Verteilungen, deren Wandel und den Möglichkeiten, diese zu verändern handelt. Damit ist die Epidemiologie die Wissenschaft von den Ursachen der Bewegung von inzidenten und prävalenten Fällen, darunter auch den Gefährdungsänderungen für das Leben und für die Gesundheit der Menschen im Ablauf der Zeit. Dieses Studieninteresse war zwar nicht neu, verlangte allerdings, die traditionelle Deutung von Epidemien als Folgen einer

Übertragung von Krankheitsursachen von Mensch zu Mensch oder Tier zu Mensch zu erweitern.

Unter dieser Voraussetzung ist für dieses Heft unseres „Kompendiums der Sozialmedizin“ die Epidemiologie die Grundlage der Analytik aller Gesundheitsprobleme, die auf Bevölkerungsebene dargestellt werden können. Sie ist Teil, der Public Health Wissenschaften, der Sozialmedizin, der Versorgungsforschung und aller weiteren Wissenschaften und Anwendungsgebieten mit analogen analytischen Bedarfen.

Ärztinnen und Ärzte, die sich speziell dem Arbeitsfeld der medizinisch gutachtlichen Expertise zuwenden wollen, werden in der Epidemiologie eine ihrer wichtigen wissenschaftlichen Grundlagen finden, hier aber auch mit den Dilemmata der Begutachtung konfrontiert sein. Diese erwachsen daraus, dass es die Qualifikation des Gutachters sein muss, probabilistische Aussagen aus Studien auf deterministisch kausale Zusammenhangsfragen bei einer Expertise zum einzelnen Fall anzuwenden. Dies wird nur gelingen, wenn reflektierte Erfahrung und die Abwägung der Möglichkeiten aber auch der Grenzen quantitativer epidemiologischer Analytik zueinander finden.

Die Herausgeber

Inhaltsverzeichnis

1	Einführung in die Allgemeine Epidemiologie.....	209
1.1	<i>Wahrscheinlichkeitskalküle in der Epidemiologie.....</i>	<i>215</i>
1.1.1	Die Quantifizierung von Erkrankungswahrscheinlichkeiten.....	216
1.1.2	Kausalität in der Epidemiologie	217
1.1.3	Die Zufälligkeit einer Verteilung.....	220
1.1.4	Wahrscheinlichkeiten für die Richtigkeit/Falschheit einer Hypothese	222
1.1.5	Ungleichheit in der Epidemiologie	223
1.1.6	Epidemiologische und klinische Studien.....	225
2	Grundbegriffe der Epidemiologie.....	229
2.1	<i>Epidemien und Regressionen.....</i>	<i>229</i>
2.2	<i>Epidemiologische Potenziale</i>	<i>232</i>
2.3	<i>Der epidemiologische Fall</i>	<i>239</i>
2.4	<i>Früherkennung</i>	<i>242</i>
2.5	<i>Epidemiologie und Ätiologieforschung.....</i>	<i>243</i>
2.6	<i>Erneuerungsmengen.....</i>	<i>248</i>
2.7	<i>Die Reproduktionsrate</i>	<i>249</i>
2.8	<i>Diagnostische und therapeutische Intervalle.....</i>	<i>250</i>
2.9	<i>Die epidemiologische Durchdringung.....</i>	<i>251</i>
2.10	<i>Die Intensität von Expositionen.....</i>	<i>252</i>
2.11	<i>Die Zeit unter Risiko.....</i>	<i>253</i>
3	Die Klassifikation von Epidemien und Regressionen.....	255
4	Quantifizierung in der Epidemiologie.....	259
4.1	<i>Das Allgemeine epidemiologische Krankheitsmodell</i>	<i>262</i>
4.2	<i>Zeitpunkte und Intervalle</i>	<i>264</i>
4.3	<i>Die epidemiologischen Studientypen</i>	<i>265</i>
4.3.1	Studien mit Primärdaten	265
4.3.2	Studien mit Sekundärdaten.....	267
4.4	<i>„Big Data“ in der epidemiologischen Forschung.....</i>	<i>268</i>
4.5	<i>Ziffern in der Epidemiologie</i>	<i>271</i>
4.5.1	Die Quantifizierung von Zuständen	272

4.5.2	Die Quantifizierung von Ereignissen.....	274
4.5.3	Die Inzidenz.....	274
4.5.4	Die kumulativen Inzidenzraten, engl. cumulative incidence rate, CI ..	275
4.5.5	Inzidenzdichte (auch Force of Morbidity, momentane Inzidenzrate, Person-Zeit-Inzidenzrate, Hazardrate)	276
4.5.6	Heilungs- oder Genesungsraten (engl. recovery rate).....	277
4.5.7	Quantifizierung von Sterbeereignissen	278
5	Relationale Häufigkeiten in der Epidemiologie	284
5.1	<i>Das relative Risiko</i>	284
5.2	<i>Odds und Odds Ratio</i>	285
5.3	<i>Das attributierbare Risiko</i>	286
6	Fehler in epidemiologischen Studien.....	288
6.1	<i>Confounder</i>	288
5.2	<i>Bias</i>	288
7	Der Zusammenhang zwischen Inzidenz und Prävalenz.....	290
8	Der Vergleich in der epidemiologischen Forschung.....	292
8.1	<i>Die Vergleichbarkeit des gemessenen Nominators</i>	294
8.2	<i>Die Vergleichbarkeit des Denominators</i>	294
9	Die epidemiologischen Bewegungsformen.....	296
9.1	<i>Systematik der Ursachen von Bewegungen der Prävalenz</i>	298
9.1.1	Demografische Faktoren.....	299
9.1.2	Medizinische Faktoren	299
9.1.3	Epidemische Faktoren.....	303
9.1.4	Präventive Faktoren.....	304
9.1.5	Interaktion der Bewegungsfaktoren	305

Abbildungsverzeichnis

Abb. 1: Allgemeine Epidemiologie und Spezielle Epidemiologie.....	210
Abb. 2: Beziehung von Falldauer, Alter und epochaler Zeit.....	224
Abb. 3: Klassifikation von Epidemien.....	255
Abb. 4: Klassifikation von Regressionen.....	257
Abb. 5: Das Allgemeine Sozialmedizinische Krankheitsmodell (Niehoff, J.-U. Der Morbiditätsprozess. Habilitationsschrift, Humboldt Universität Berlin 1983)	263
Abb. 6: Diagnostische und therapeutische Intervalle	264
Abb. 7: Die Erneuerung der Prävalenz ((Niehoff, J.U., Li, B. Epidemics and Regressions. Saluscon Akademie Verlag, Berlin 2022, ISBN 978-3-948267-15-5) .	290
Abb. 8: Die Bewegungsformen der Prävalenz und ihre Ursachen (nach B. Kreuz und Mitarbeiter: Über den Zusammenhang zwischen Krankheitsdauer und Morbidität. Ztschr. ärztl. Ausbildung 67 (1973) 144-148).....	296
Abb. 9: Der Zusammenhang von Normwertänderungen und diagnostizierten Fällen	301

Tabellenverzeichnis

Tab. 1: Vierfeldertafel für den Zusammenhang von Krankheit und Testresultat	241
Tab. 2: Gütekriterien eines diagnostischen Tests	242
Tab. 3: Zusammenhang von Prädiktion (PV), Likelihood-Ratio (LR) und Prävalenz (nach B. Höldke).....	243
Tab. 4: Klassifikation von Epidemien	256
Tab. 5: Klassifikation von Regressionen	258
Tab. 6: Beispiele für epidemiologisch relevante Ereignisse, Zustände und Intervalle	261
Tab. 7: Sterbefälle (alle Ursachen), differenziert nach dem Umgang mit Alkohol, modif. nach: Miller, G.J. et al, in: Inter. J. of Epidemiology, 19:1990, p. 928	285
Tab. 8: Relatives und attributierbares Mortalitätsrisiko für Lungenkrebs und koronare Herzkrankheit bei Rauchern in einer Kohortenstudie unter britischen Ärzten, jährliche Mortalitätsrate pro 100.000 (zitiert nach: R. Doll and R. Peto, Mortality in relation t..	287
Tab. 9: Simulation der Wirkungen von Zeitverschiebungen auf die Struktur der Fälle nach dem Merkmal "Krankheitsstadium"	302

1 Einführung in die Allgemeine Epidemiologie

Die Epidemiologie ist die Wissenschaft vom Wandel der Gesundheitsprobleme definierter Bevölkerungen. Die Formen dieses Wandels sind Epidemien und Regressionen und deren Bilanz bestimmt die Richtung der epidemiologischen Transition.

Auf dem Fundament der „Allgemeinen Epidemiologie“ haben sich vielfältige Anwendungsgebiete entwickelt. Zu diesen gehört auch die Nutzung durch die Sozialmedizin und ihr Anwendungsfeld „Public Health“.

Die Verteilungen gesundheitlicher Belastungen in einer Bevölkerung können zufällig (umgangssprachlich auch „schicksalhaft“), oder aus nachweisbaren Ursachen ungleich verteilt sein. Die Beobachtungen der Zunahmen und der Verringerungen von Gefährdungen für die Gesundheit im Ablauf der Zeit gehen regelhaft mit der Zunahme oder der Abnahme von Gefährdungsunterschieden zwischen Teilen der Bevölkerungen einher. Veränderungen der Inzidenzen und Prävalenzen, also der Mengen (im Englischen auch „Frequenzen“) neuer Erkrankungen und der Mengen der Fälle im Zustand des Krankseins, sind somit die sekundären Folgen von primären Prozessen des Wandels von Gefährdungsverteilungen. Diese Wandlungen nennen wir Bewegungen und sie sind es, die Ungleichverteilungen bewirken. Solche Bewegungen sind keine Eigenschaften von Krankheiten. Krankheiten können weder zu- noch können sie abnehmen. Zu- oder abnehmen können allein die Mengen der von solchen Bewegungen Betroffenen und die diesen zuzuordnenden „epidemiologischen Fälle“.

Die Identifizierung von „epidemiologischen“ Bewegungen verlangt zunächst die Charakterisierung der Teile einer Bevölkerung, die solche Veränderungen erleben. Von den vielfältigen Bewegungsursachen finden Zunahmen wie auch Minderungen von Gefährdungen und die Folgen für die davon betroffenen Teile einer Bevölkerung stets besondere Aufmerksamkeit. Sie repräsentieren dann jeweils Vor- oder Nachteile für einzelne Cluster von Bevölkerung und werden dann in der öffentlichen Wahrnehmung als Ungleichheit, ggf. als Benachteiligung interpretiert.

Erkenntnisse zu solchen ungleichen Verteilungen sind die Wissensbasis der sozialen Prävention, dienen der Bedarfs- und Versorgungsforschung sowie der Folgeabschätzung gewollter wie ungewollter Einflussnahmen auf die Gesundheitsverhältnisse. Solche Einflussnahmen sind vielschichtig und haben zur Entstehung einer Vielzahl von speziellen Anwendungsgebieten der Epidemiologie geführt. Allen gemeinsam sind die methodischen Prinzipien der Beschreibung und der Analyse entsprechender Phänomene.

Aus systematischen Gründen ist es sinnvoll, das Arbeitsfeld der Epidemiologie hinsichtlich ihrer allgemeinen Prinzipien (Allgemeine Epidemiologie) und ihrer spezifischen Arbeitsfelder (spezielle Epidemiologie, siehe Heft 6, 7, 8) zu unterscheiden.

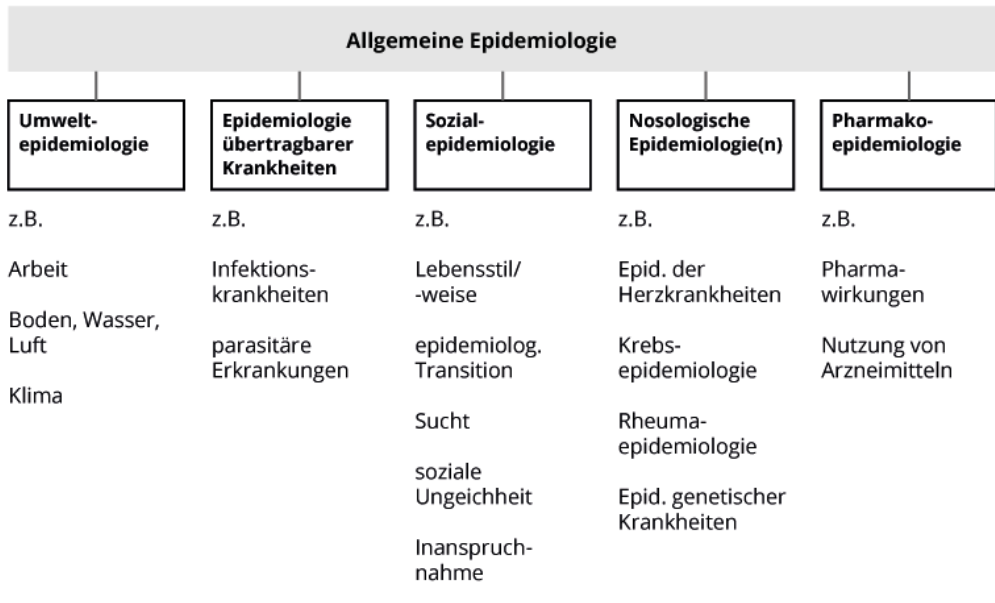


Abb. 1: Allgemeine Epidemiologie und Spezielle Epidemiologie.

Die Epidemiologie

- beschreibt die gesundheitlich relevanten Eigenschaften von Bevölkerungen (z.B. Art, Häufigkeit, Verteilung oder Folgen von Expositionen, von Krankheiten, von Behinderungen oder von Pflegebedarfen) – Phänomenologie
- untersucht die Ursachen der Veränderungen in der Verteilung solcher Eigenschaften – Kausalität
- bewertet die Folgen solcher Veränderungen für die Prävention und die Gesundheitsversorgung – Assessment von Ungleichheiten

Epidemiologische Analysen müssen aus mindestens drei Gründen vom Individuum abstrahieren:

1. Die Häufigkeit „1“ ist epidemiologisch irrelevant. Ein einzelnes Individuum oder ein einzelner epidemiologischer Fall können keine Objekte einer quantitativen epidemiologischen Analyse sein.
2. Der Einzelfall erlaubt keine vergleichenden Aussagen über Gefährdungsänderungen und damit auch keine auf Bevölkerungen oder ihre Gliederungen zu beziehenden Handlungsorientierungen.
3. Die Ursache eines einzelnen Falls kann nur mit deterministischen ex post Analysen geklärt werden. Epidemiologisch kausale Analysen sind probabilistischer Natur und nur für ex ante, bzw. prädiktive Feststellungen und Handlungsziele nutzbar.

Die medizinischen Wissenschaften suchen gemeinsam und arbeitsteilig Antworten u.a. auf vier Fragen:

1. Warum werden Menschen krank? (Ätiologie)
2. Wie ist der Verlauf einer Krankheit zu erklären? (Pathogenese)
3. Wie und warum verändert sich die Anzahl der an einer Krankheit leidenden Menschen im Ablauf der Zeit? (Ursachen von zunehmenden Prävalenzen (Epidemien) oder von abnehmenden Prävalenzen (Regressionen)).
4. Warum ist die Anzahl der an einer Krankheit erkrankenden Personen in einer Bevölkerung nicht zufällig verteilt? (Ursachen und Auswirkungen von Ungleichheit)

Die letzten beiden Fragen betreffen den Kern des epidemiologischen Forschungsfeldes. Die untersuchungsrelevanten Phänomene lassen sich dann allgemein unter den Begriffen Ereignis (z.B. Krankwerden), Zustand (z.B. Kranksein) und Zustandsdauer (z.B. Lebensdauer unter Risiko, Krankheits- und Behandlungsdauer etc.) beschreiben. Hierfür sind jeweils geeignete Indikatoren und deren Bezug zur räumlich und zeitlich zugehörigen Grundgesamtheiten (Bevölkerungen, epidemiologische Potenziale), bzw. zur Lebensdauer einer Bevölkerung oder einer Teilbevölkerung (z.B. Lebensdauer unter Risiko) erforderlich.

Die Ursachen der Zunahme (Epidemien) und der Abnahme (Regressionen) von Gefährdeten und von kranken Personen sowie die Ursachen der nicht zufälligen und ungleichen Verteilung von Krankernden und Kranksein in einer Bevölkerung sind das Forschungsfeld der Epidemiologie.

Dieser konzeptionelle Ansatz schließt alle Expositionen, Krankheiten, Behinderungen und sonstige gegeneinander abgrenzbaren Gesundheitsprobleme sowie die von ihnen Betroffenen ein.

Unter der Nullhypothese, dass die Varianz der biologischen Eigenschaften in einer Bevölkerung einer Zufallsverteilung folgt und in aufeinander folgenden Geburtsjahrgängen hinreichend konstant reproduziert wird, bedürfen alle Abweichungen von dieser Hypothese einer ursächlichen Erklärung. Sie wären dann allgemein die Folge von Änderungen (Bewegungen) der relevanten gesundheitsbezogenen Einflüsse auf Bevölkerung.

Die Ursachenanalyse von Verteilungsänderungen und -unterschieden ist folglich eine Analyse der Ursachen unterschiedlicher Bewegungen von gesundheitsrelevanten Einwirkungen auf Teile einer Bevölkerung. Diese führen dann zu Verteilungsunterschieden, die durch den Zufall nicht erklärt werden können. Solche Bewegungen können

mittels ihrer Richtungen und ihrer Dynamik sowie mittels ihrer Intensitäten und ihrer Dauer beschrieben werden.

Beispiel:

Die Menge neuer Fälle der Krankheit A betrug zu Beginn eines Beobachtungsintervalls 100 und an dessen Ende 150. Es hat also eine Bewegung der neuen Fallmengen mit der absoluten Größe 50 gegeben. Das populäre Interesse an der Deutung dieser Beobachtung bezieht sich auf die qualitative Bewertung dieser Veränderung als „gut“ oder „schlecht“. Die Richtung dieser Bewegung hat in diesem Beispiel ein positives Vorzeichen. Über die Zeit können solche Bewegungen in ihrer Intensität (synonym Stärke, Kraft, Dichte) variieren. Das gilt analog auch für die Richtung. Schließlich ist die Bewegung auch über ihre Dynamik oder Geschwindigkeit zu beschreiben. Der umgangssprachliche Begriff für Dynamik ist das Tempo. Die Mengenänderung von A ist also durch ihre Richtung, ihre Intensität und ihr Tempo vollständig beschrieben. Diese Bewegungen können in der Bevölkerung einheitlich, in Teilmengen (Subgruppen) aber auch unterschiedlich sein. Letzteres erzeugt dann Unterschiede zwischen den Teilbevölkerungen. Die Unterschiede der Mengen Gefährdeter und Kranker in einer Bevölkerung sind also die Folge von Bewegungen. Dementsprechend sind Bewegungen und deren Ursachen das Schlüsselthema der Epidemiologie.

Richtung, Intensität und Tempo dienen in der Epidemiologie zur Beschreibung von Bewegungen der Prävalenz.

Für das Verständnis epidemiologischer Bewegungen sind weitere Sachverhalte näher zu bestimmen:

1. Zeit

Die Variable „Zeit“ ist für die Erklärung epidemiologischer Phänomene von grundlegender Bedeutung. Ohne eine zeitliche Zuordnung können Bewegungen nicht beschrieben werden. Eine solche Zuordnung ist in der Praxis häufig kein triviales Problem. Menschen betrifft die „Zeit“ in drei Kontexten: Ontogenese, Epoche und Biografie. Während das Lexis-Diagramm (siehe Heft 2) die Beziehungen von Epoche und kalendarischem Alter grundsätzlich ordnet, kann der Umgang mit dem ontogenetischen oder biologischen Alter und mit der individuellen Biografie eine erhebliche Herausforderung sein. Alterung und soziale Biografien sind interpersonal variant und intrapersonell modifikabel. Gleiche Beobachtungen können also sehr unterschiedlichen zeitlichen Bezügen zuzuordnen sein.

Kalendarisch Gleichaltrige sind in der Regel biologisch unterschiedlich alt. Vergleiche von Teilbevölkerungen sind deshalb immer dann schwierig, wenn die Variable Alter (als Maß für Zeit) die untersuchten Sachverhalte, z.B. altersspezifische Häufigkeiten, beeinflusst. Das kann bei der Bewertung entsprechender Studienergebnisse in methodische Herausforderungen führen. Analog gilt dies auch für das Verhältnis von epochaler und ontogenetischer Zeit. Da das biologische Altern modifikabel ist, können Merkmalsverteilungen gleicher Sachverhalte sich in der Generationenfolge auf unterschiedliche biologische Lebensalter beziehen. So kann etwa die Forderung nach Altersgleichheit zweier zu vergleichender Studienbevölkerungen auch dann nicht erfüllt sein, wenn die kalendarische Altersgleichheit gewährleistet ist. Dieser Effekt verstärkt sich bei globalen und epochalen Vergleichen und hat in unterschiedlichen Lebensaltern ein unterschiedliches Gewicht.

2. Krankwerden

Krankwerden ist ein Ereignis, das z.B. den Wechsel vom Zustand „Gesundsein“ in den Zustand „Kranksein“ bewirkt. Epidemiologische Ereignisse sind also grundsätzlich Zustandswechsel zu Zeitpunkten. Allerdings lassen sich diese Zeitpunkte bei vielen Krankheiten nicht exakt ermitteln und auch nicht gut schätzen. Zwischen den Ereignissen des Krankwerdens, ersten Arztkontakten und einer korrekten Diagnosestellung können (erhebliche) Zeitspannen liegen, die bei verschiedenen Krankheiten, aber auch bei gleichen Krankheiten in unterschiedlichen Bevölkerungen und deren Gliederungen deutlich variieren können.

Es ist zudem eher selten, dass alle Ereignisse des Krankwerdens in einer Bevölkerung auch bekannt werden. Je katastrophaler solche Ereignisse erlebt werden und je gesicherter der Zugang zum medizinischen Versorgungssystem ist, desto eher kann davon ausgegangen werden, dass über die Diagnosestellungen die Ereignisse des „Krankwerdens“ weitgehend vollständig und mit kurzen diagnostischen Intervallen abgebildet werden. In diesem Fall kann die Häufigkeit einer Diagnosestellung als Schätzer für die Häufigkeit des Krankwerdens gelten. Dies gibt jedoch keine Gewähr, dass die dokumentierten Mengen vollständig sind.

Es ist zusätzlich zu beachten, dass bei wiederholbaren Krankheiten die Anzahl der Fälle von Krankwerden in einem Beobachtungszeitraum die Zahl der jeweils neu erkrankenden Personen deutlich übersteigen kann. Für epidemiologische Fälle gibt es jedoch nur dann einen sinnhaften Bezug zur Bevölkerung, wenn diese nach der Ereignisfolge unterschieden werden können. Dies ist notwendig, weil bei einer Reihe von Krankheiten eine Person innerhalb eines Beobachtungszeitraums mehr als einmal erkranken kann. Mit anderen Worten: Die Ermittlung „echter“ Wahrscheinlichkeiten für das Krankwerden ist für viele der vorkommenden Krankheiten ohne präzises Studiendesign nicht möglich.

3. Kranksein

Kranksein ist der Zustand zwischen den Ereignissen Krankwerden und Ausscheiden aus der Menge der Kranken. Ein solcher Zustand ist über die Krankheitsdauer definiert. Im Regelfall bleibt die Krankheitsdauer jedoch unbekannt. Etwas leichter ist die Behandlungsdauer zu messen. Bei einem kurzen diagnostischen Intervall ist die Behandlungsdauer ein hinreichend guter Schätzer der Krankheitsdauer. Dann kann aus Änderungen der Behandlungsdauer auf Änderungen der Krankheitsdauer in einer Bevölkerung geschlossen werden.

4. Empfänglichkeit oder Disposition

Werden z. B. 1.000 Personen einer exakt gleichen schädigenden (biologischen, chemischen, physikalischen oder psychischen) Einwirkung ausgesetzt, werden die Folgen variieren. Menschen sind unterschiedlich disponiert oder anfällig gegenüber einer Gefährdung (individuelle Suszeptibilität). Über die Ursachen dieser Variationen ist, zumindest auf Bevölkerungsebene, noch wenig bekannt. Konzepte und Begriffe wie genetische Disposition oder Prädisposition, sensitive Lebensabschnitte, Altersspezifik, Fitness, Anfälligkeit, Widerstandsfähigkeit, Anpassung, Gewöhnung, Konstitution oder Arbeits- und Lebensverhältnisse sind oft kaum belastbar definiert. Dispositionen sind jedoch ein entscheidender Grund der Heterogenität und der Nichtzufälligkeit in der Verteilung der in einer Bevölkerung vorkommenden Krankheiten.

Das gilt auch in der Umkehrung: Werden 1.000 Personen mit exakt gleichen Wirkungen (hier Krankheiten) untersucht, können diese dennoch Folgen sehr unterschiedlicher Ursachen sein. Der hierfür oft benutzte Begriff der Multikausalität ist auf den einzelnen Fall nicht anwendbar, macht also nur in Bezug auf die Gesamtzahl der ursächlich heterogenen Fälle Sinn.

5. Konkurrierende Gefährdungen, konkurrierende Risiken

Auf der Bevölkerungsebene können Gefährdungen und ihre Auswirkungen miteinander konkurrieren. Wenn es das Ziel ist, eine bestimmte Gefährdung für die Gesundheit auszuschließen, bzw. sie zu vermeiden, kann dies im Erfolgsfall bewirken, dass die Übergangswahrscheinlichkeiten in andere wirksame Gefährdungsfolgen positiv oder negativ beeinflusst werden. Die erfolgreiche Eingrenzung etwa der Tuberkulose kann also zur Folge haben, dass „konkurrierend“ andere Gefährdungen, z. B. Karzinom-erkrankungen davon „profitieren“. Der „Gewinn“ besteht dann ggf. darin, dass die Erkrankungsalter zunehmen, es also einen Gewinn an Lebenszeit ohne das Leben bisher bedrohender Krankheiten gibt.

Epidemiologische Ereignisanalysen haben folglich zu prüfen, ob und welche Übergänge in andere Gefährdungen des Krankwerdens anzunehmen sind, als Möglichkeit also gleichsam durch die Vermeidung bestimmter Gefährdungen neu geschaffen werden.

1.1 Wahrscheinlichkeitskalküle in der Epidemiologie

Wahrscheinlichkeiten werden umgangssprachlich als Risiken und als Chancen bewertet.

Wahrscheinlichkeiten können im Kontext epidemiologischer Studien quantifiziert werden, wenn hierfür die methodischen Voraussetzungen erfüllt sind. Diese sind die Kenntnis der Ausgangsmengen der Personen für die entsprechende Kalküle vorgenommen werden sollen, die Kenntnis der Einwirkungen, die den Chancen oder den Risiken zugerechnet werden und die Zeitdauer, für die die Kalküle gelten sollen. Diese haben dann Wertebereiche zwischen 0 und 1. Für die Unterscheidung von Chance und Risiko gibt es hingegen kein Maß. Es handelt sich um subjektive Wertungen oder soziale Präferenzen, die den Aufgaben der Epidemiologie kaum zugerechnet werden können.

Da für nahezu beliebige disjunkte Ausgangsmengen von Personen solche Kalkulationen vorgenommen werden können, sind entsprechend beliebig viele Wahrscheinlichkeitskalküle und -unterscheidungen möglich. Die Epidemiologie hat es immer mit inhomogenen Mengen zu tun, also mit Personen, die mittels ihrer Strukturvariablen zu unterscheiden sind. Für die Epidemiologie sind diese „Mengen“ immer Bevölkerungen oder Teilbevölkerungen, die hier als „epidemiologische Potenziale“ bezeichnet werden. Analog gebräuchliche Begriffe sind „Cluster“, „Sample“, „Risikogruppen“ und in klinischen Settings auch „Patientengruppen“ oder „Genesene“ etc. Es handelt sich also immer um Teilmengen, für die jeweils deren Risiken oder Chancen zu messen sind.

Es ist naheliegend zu argumentieren, dass die Menge neuer Erkrankungen (I) dem Produkt der für eine Bevölkerung oder Teilbevölkerung wirksamen a priori Wahrscheinlichkeit (r) zu erkranken und dem Umfang dieser Bevölkerung unter Risiko (N) direkt proportional ist:

$$I_{(x)} = r \times N$$

(gültig unter Gleichgewichtsbedingungen von Zu- und Abgängen der Bevölkerung unter Risiko).

Das gilt analog für alle Potenzialmengen, also auch für Gefährdete und Nichtgefährdete, für Kranke und für Personen, die an den Folgen von Krankheiten leiden. Sowohl Änderungen von „r“ also auch Änderungen von „N“ können zu einer Änderung der absoluten Mengen von „I“ führen. Um zu verstehen, weshalb sich die Menge von Neuerkrankungen in einem Beobachtungszeitraum ändert, sind folglich die Größenbewegungen von r und/oder von N zu ermitteln. Der Begriff „epidemiologisches Potenzial“ schließt ein, dass Bevölkerungen stets inhomogen sind, also aus einer Vielzahl von

Potenzialen bestehen, die dann die Struktur einer Bevölkerung abbilden. Hieraus folgt dann auch, dass die Epidemiologie die Struktur von Bevölkerungen zu ihrem Forschungsobjekt hat und die jeweils relevanten Einwirkungen auf diese Bevölkerungen deren Forschungsgegenstände sind. Letztere beziehen sich auf alle gesundheitsrelevanten Einwirkungen, also sowohl auf positive (präventive) wie negative (gefährdende) Einwirkungen. Und sie beziehen sich auch grundsätzlich auf alle Einwirkungen und auf alle Krankheiten, die nicht einer Zufallsverteilung folgen. Sie betreffen also übertragbare und nicht übertragbare Krankheiten und Gefährdungen sowie Gewalt- und Katastrophenfolgen in gleicher Weise.

Wenn die jeweiligen Einwirkungen (Expositionen, Impakts) nur Teile einer Bevölkerung betreffen, folgen ihnen immer auch Unterschiede in der Menge der Neuerkrankungen in diesen Teilbevölkerungen. Folglich ist die Ursächlichkeit von unterschiedlich häufigen Neuerkrankungen in Teilen der Bevölkerung erklärt, wenn die unterschiedlichen Dynamiken von „r“ und/oder „N“ in jedem dieser Potenziale erklärt sind.

1.1.1 Die Quantifizierung von Erkrankungswahrscheinlichkeiten

Die Messung von Erkrankungswahrscheinlichkeiten verlangt die Zählung der Ereignisse des Krankwerdens an einer definierten Krankheit und deren Bezug auf das zugehörige epidemiologische Potenzial.

Um alle bei Menschen möglichen Krankheiten zu erfahren, muss eine Bevölkerung maximal groß und heterogen sein sowie über die gesamte Dauer ihrer Existenz beobachtet werden (können). Diese Voraussetzung ist nicht zu erfüllen. Messungen der Epidemiologie sind folglich immer zeitlich und räumlich begrenzt und solche Messungen sind diesen Grenzen auch zuzuordnen. Daraus folgt, dass sich die Häufigkeiten des Krankwerdens in Zeit und Raum bewegen (verändern) und sich im Vergleich der epidemiologischen Potenziale unterscheiden oder auch unverändert bleiben können. Struktur und Prozess sind somit die Schlüsselphänomene epidemiologischer Forschung, die der Erklärung bedürfen.

Wenn es die Möglichkeit gibt, zu erkranken oder nicht zu erkranken, sind diese Alternativen durch ihre Ereigniswahrscheinlichkeiten zu beschreiben. Es ist folglich zu prüfen, ob und unter welchen Voraussetzungen in der Epidemiologie Wahrscheinlichkeiten gemessen werden können.

Eine erste Voraussetzung ist dann erfüllt, wenn die Anzahl neuer Ereignisse des Krankwerdens in einem Zeitraum bekannt ist. Dies ist der Fall, wenn die Menge der dokumentierten richtigen Diagnosestellungen der Menge der neuen Erkrankungen (hinreichend) entspricht. Diese Forderung kann bei einigen (wenigen) Krankheiten gut erfüllt werden, bei vielen anderen nicht. Allerdings ist sicherzustellen, dass die Sensitivitäten und die Spezifitäten der diagnostischen Tests bekannt sind und diese in ihrer Anwendung und bewertenden Funktion standardisiert sind.

Schwieriger ist der Bezug dieser Messungen auf die Gesamtheit der möglichen Ereignisse, also auf die epidemiologischen Potenziale. Die etwa an ein Lottospiel geknüpfte Forderung, dass jede Kugel die gleiche Chance (a priori Wahrscheinlichkeit) haben muss, gezogen zu werden, kann auf eine Bevölkerung nicht angewandt werden. Die Wahrscheinlichkeit krank sowie dann auch diagnostiziert und behandelt zu werden, sind tatsächlich vielfältig bedingt (bedingte Wahrscheinlichkeit) und in Bezug auf einzelne Gefährdungen und deren potenzielle Folgen hochgradig verschieden. Es gibt jenseits spekulativer Erwägungen praktisch nie Grund zu der Annahme, die Vulnerabilitäten gegenüber einer Exposition oder deren Durchdringung einer Bevölkerung seien gleichverteilt.

Für Bevölkerungen *a priori* Wahrscheinlichkeiten für alle vorkommenden Krankheiten zu bestimmen, verlangte die Annahme der Homogenität einer Bevölkerung bezüglich dieser Möglichkeiten. Das gilt entsprechend, wenn nicht die Bevölkerung als Bezug gewählt wird, sondern die Dauer, die diese einer Gefährdung ausgesetzt ist. Das Homogenitätsgebot würde also auch für die „Zeit unter Risiko“ gelten. Solche Homogenitätsannahmen sind zumeist realitätsfremd. Die Personen einer Bevölkerung sind von der Gesamtheit der möglichen Krankheitsgefährdungen nicht in gleicher Weise betroffen. Es gibt Personen, die aus welchen Gründen immer, für bestimmte Krankheiten prädisponierter sind als andere. Es ist auszuschließen, dass eine Person, bzw. eine Gruppe einander ähnlicher Personen alle bekannten Krankheiten je bekommen könnte. Es ist für viele potenzielle Krankheitsursachen bekannt, welche Wirkungen sie haben können. Weitgehend unbekannt ist hingegen, warum manche Menschen trotz intensiver Exposition nicht erkranken. Diese Frage wird interessanter Weise auch nur selten gestellt.

Es müsste für die Bestimmung von *a priori* Wahrscheinlichkeiten folglich gelingen, Teilbevolkerungen zu erzeugen, die homogen sind. Dies kann nur unter Laborbedingungen gelingen und auch dann bestenfalls in Annäherung. Aus Sicht der Epidemiologie hieße dies dann aber, das wichtigste und im Besonderen für die medizinische Praxis herausragende Merkmal von Bevölkerungen, nämlich deren Heterogenität, zu eliminieren.

1.1.2 Kausalität in der Epidemiologie

Ursachen erklären Wirkungen. Ursachen und Wirkungen sind folglich immer aufeinander bezogen.

Die möglichst exakte Benennung einer Wirkung ist eine notwendige Bedingung dafür, etwas zeitlich Vorausgehendes als Ursache bezeichnen zu können. Je komplexer die benannte Wirkung, desto uneindeutiger ist dann auch die Ursache, bzw. sind die benennbaren Ursachen. Der Begriff der Multikausalität steht deshalb vor allem für die mangelnde Präzision bei der Definition einer zu erklärenden Wirkung.

Die Abgrenzung einer zu erklärenden Wirkung bestimmt folglich die Präzision, mit der von einer Ursache gesprochen werden kann. Die primäre Frage in der Ursachenanalytik ist also stets, welche spezifische Wirkung kausal erklärt werden soll. Die Frage, warum Patient A an einer Glomerulonephritis erkrankt ist, kann eine epidemiologische Studie nicht beantworten. Deren Aufgabe wäre es hingegen, zu untersuchen, warum in einer definierten Bevölkerungsgruppe und in einem definierten Zeitraum die Zahl der neuen Erkrankungsfälle an einer Glomerulonephritis zu- oder abgenommen hat bzw. in einer Bevölkerung ungleich verteilt ist.

In der Epidemiologie, anders als unter klinischen oder unter experimentellen Studienbedingungen, ist die exakte Gleichartigkeit von zu untersuchenden „Wirkungen“ zu meist nicht zu sichern, ggf. auch nicht gewollt. Zustands- wie Ereignishäufigkeiten bei als gleich bezeichneten Krankheiten sind bei bevölkerungsbezogenen Vergleichen regelhaft variant. Weder die „Gleichheit“ der Untersuchungsgegenstände noch die der Untersuchungsobjekte wird sich wirklich sichern lassen. Das hat für die Frage nach der „Kausalität“ im Falle epidemiologischer Phänomene Folgen: Die Ätiologie eines einzelnen individuellen Krankheitsfalls kann durch die Epidemiologie nicht geklärt werden, da gleiche Erkrankungsfälle in ihren Ursachen variieren können.

Beispiel:

Die Erfahrung, dass quantitativ begründete ursächliche Aussagen auf einen Einzelfall ggf. nicht zu übertragen sind, hat in der medizinischen Erkenntnisgeschichte vielfältig eine Rolle gespielt. Sie führte z.B. den Physiologen Max Verworn (1892-1921) und den Pathologen David von Hansemann (1858-1920) zu der Überzeugung, Krankheiten ließen sich ursächlich nicht erklären, da ihre Entstehung in sehr komplexe Bedingungen eingebettet sei. Das Konzept des „Konditionalismus“ sollte deshalb das der „Kausalität“ als wissenschaftliches Leitkonzept ersetzen. Die Vorstellungen der Vertreter der Sozialhygiene/Sozialmedizin der 1920er Jahre waren hierzu weitgehend konform. Die „Soziale Pathologie“ von A. Grotjahn (1869-1931) (Soziale Pathologie“; Springer Verlag, Reprint 1977, ISBN-13: 978-3-642-93040-9 DOI: 10.1007/978-3-642-93039-3) war entsprechend eine „soziale Bedingungslehre“ zur Entstehung von Krankheiten. Diese sollten durch den Wandel der „Bedingungen“ zur Besserung der gesundheitlichen Verhältnisse führen. Die Kausalitätssuche ist vor allem ein Erkenntnisprinzip. Handeln ist hingegen immer in komplexe Bedingungen eingebettet und hat diese auch zu beachten.

Ähnliche Erwägungen spielten in den 1970er Jahren erneut eine erhebliche Rolle. Ihr Kern bestand in der Auffassung, die Epidemiologie könne durch die Ermittlung der gesamten als Risiken zu interpretierenden Bedingungen für die Entstehung von Krankheiten auch ohne weitere Kenntnisse der Kausalität die Voraussetzungen für die sukzessive Eradikation der wichtigsten Krankheiten, vor allem der Krankheiten des Herzens und des Kreislaufs, der Stoffwechselkrankheiten und der Karzinome

schaffen. Hierzu bedürfe es allein eines Lebens, das pathogene Bedingungen überwindet und sie durch solche ersetzt, die den Bedingungen entsprechen, unter denen diese Krankheiten selten sind. Hierbei gab es jedoch auch einen entscheidenden Perspektivwandel: die komplexen sozialen Lebensbedingungen wurden durch individuelle und als selbstverantwortete Entscheidungen über das individuelle Verhalten ersetzt, also gleichsam desozialisiert, ein Wandel, der den Lebenssichten von Mittel- und Oberschichten entgegenkam. An die Stelle einer Epidemiologie, die sich mit der gesellschaftlichen Lebensweise auseinandersetzte, trat eine Epidemiologie der „Lebensstile“. Die „klassischen Bedingungsthemen“ „Arbeit“, „soziale Chancenverteilung“ und „Selbstbestimmbarkeit“ wurden in erheblichem Umfang durch die Orientierung auf individuelle Lebensstile ersetzt. Folgerichtig wurden dann deren gesellschaftliche Tolerierbarkeit und das Argument des Selbstverschuldens von Krankheit zum globalen Leitthema der öffentlichen Debatten um prioritäre Gesundheitsprobleme und zu den Sicherungssystemen im Krankheitsfall.

Kausalitätsforschung durch die Epidemiologie zielt auf zwei Schlüsselthemen. Das sind

- die Ungleichverteilungen des Krankwerdens und des Krankseins in einer Bevölkerung und
- die Veränderung solcher Verteilungen im Ablauf der Zeit

Mit jedem Erklärungsansatz zu den entsprechenden Phänomenen verbinden sich auch eigene Handlungsorientierungen für die Prävention, für die Zugangssicherung zu medizinischen Hilfen und für die sozialen Konzepte zur Bewältigung von Krankheitsfolgen.

Veränderungen sind im Kontext der Epidemiologie Bewegungen von einem Zustand in einen anderen. Solche „Bewegungen“ lassen sich durch ihre

1. Richtung (Epidemien, Regressionen),
2. Intensität (Stärke, Kraft),
3. Geschwindigkeit (Veränderung pro Zeitintervall) und
4. Durchdringung einer Bevölkerung und ihrer spezifischen epidemiologischen Potenziale beschreiben.

Wenn solche Bewegungen die Gesamtheit der Menschen nicht gleichzeitig, gleichsinnig oder gleich stark betreffen, führen sie zu neuen Verteilungen. Sie führen auch zu ungleichen Verteilungen und Häufigkeiten in verschiedenen Teilen der Bevölkerung. Beobachtete Unterschiede in den Gesundheitsproblemen zwischen mindestens zwei Bevölkerungen sind im Wahrnehmungsbereich der Epidemiologie also die Folge (Wirkung) von ursächlichen Einwirkungen unterschiedlicher Intensität auf mindestens eine der Teilbevölkerungen in einer bestimmten Zeitspanne. Solche Bewegungen und deren quantitative und strukturelle Folgen.

Beispiel:

Mit der Entdeckung des Mycobacterium tuberculosis hatte Robert Koch (1843-1910) die Ursache der Tuberkulose beschrieben und damit der klinischen Forschung neue Wege eröffnet. Lange vor Koch war durch epidemiologische Studien aber bereits gezeigt worden, dass diese Krankheit in sozial vertikal gegliederten Bevölkerungen unterschiedlich häufig vorkam, die Menschen also in Bezug auf die Tuberkulose unterschiedlichen Gefährdungen, Empfänglichkeiten und Überlebenschancen ausgesetzt waren. Es konnte auch beobachtet werden, dass die Neuerkrankungen an Tuberkulose mit der Verbesserung der Lebensverhältnisse zurückgingen, ebenso wie sie in Zeiten der Verschlechterung der Lebensverhältnisse häufiger wurden. Jede dieser Wahrnehmungen ist eigenständiger und von der Koch'schen Entdeckung unabhängiger ursächlicher Erklärung bedürftig, ebenso wie aus solchen Erkenntnissen dann spezifische praktische Konsequenzen folgen.

1.1.3 Die Zufälligkeit einer Verteilung

Epidemiologische Merkmale liegen als diskrete oder stetige Verteilungen vor.

Merkmalsverteilungen resultieren aus den Wahrscheinlichkeiten, mit denen die Messwerte einer Variablen ihren Skalierungen empirisch zugeordnet sind. Die Anzahl solcher Messwerte je Skalenpunkt oder -klasse wird also immer von der Wahrscheinlichkeit ihres Auftretens bestimmt. In diesem Sinne repräsentieren die gemessenen Werte insgesamt die zugrunde liegende Wahrscheinlichkeitsverteilung. Die dokumentierten Werte sind Häufigkeiten, die vom Produkt der Wahrscheinlichkeit und der Klassenbesetzung durch die Elemente auf die die Wahrscheinlichkeiten wirken, bestimmt sind.

Im Falle epidemiologischer Analysen sind entsprechende Wahrscheinlichkeitsverteilungen typischer Weise nicht normal verteilt. Es handelt sich vielmehr um schiefe oder auch um mehrgipflige Verteilungen. Die Verteilung relevanter Merkmale wie Körperhöhe, Körpergewicht oder z.B. der Blutdruck sind also in einer Gesamtbevölkerung oder in einer diese Bevölkerung repräsentierenden Stichprobe nicht zufällig verteilt, bzw. weichen von einer symmetrischen (Gauß'schen) Zufallsverteilung (auch Normalverteilung) ab. Obwohl in der Epidemiologie oft mit „Durchschnitten“ argumentiert wird, sind diese tatsächlich nur Verteilungsparameter oder Kennziffern, die ohne Kenntnis der Verteilungen, die sie beschreiben, nicht problemfrei zu interpretieren sind.

Solche Kennzahlen beschreiben z.B. die mittleren Lagen einer Verteilung, ihre Varianz, bzw. die Streuung der Merkmalsverteilungen, deren Schiefe oder Asymmetrie oder deren eingipfligen oder mehrgipfligen Gestalt. Übliche Parameter sind neben den arithmetischen Mittelwerten, die Medianwerte und die Dichtemittel, aber auch

Quantile, Dezile etc. Nutzer epidemiologischer Studien, z.B. medizinische Gutachter, sollten die Interpretationsmöglichkeiten und -grenzen solcher Kennzahlen kennen.

Die Verteilung von Merkmalen, die Normalverteilungen folgen und eine konstante Varianz haben, bedürfen keiner Erklärung. Hieraus folgt, dass die Epidemiologie nicht zufällige Merkmalsverteilungen und deren Varianzen untersucht. Sie folgt dabei der These, dass Verteilungen durch spezifische äußere Einwirkungen beeinflusst und geformt sein können. Diese sind letztlich Einwirkungen aus den jeweiligen und vielfältig zu differenzierenden Lebensumwelten von Menschen. Soweit sie direkt und indirekt als von Menschen gemacht gelten können, sind sie dann auch als soziale Einwirkungen zu klassifizieren.

Neben oder zusätzlich zur genetischen Varianz gibt es also lebensumweltliche Einflüsse, welche die Vielfalt, bzw. die Inhomogenität von Bevölkerungen ursächlich erklären. Die für die Epidemiologie bedeutsamen Variablen sind also solche, die die Vielfalt der Grundgesamtheit der Elemente einer Bevölkerung, also deren Ungleichheit abbilden. Dieser Sachverhalt ist von einer Bewertung als Benachteiligung strikt zu unterscheiden. Die für die Epidemiologie relevanten Merkmale einer Bevölkerung, hier die Merkmale, die gesundheits- und krankheitsassoziiert sind, folgen also keinen symmetrischen Verteilungen, und sie reproduzieren sich in ihren zeitlichen Parametern (Epoche, Generation, Lebensalter) typischer Weise auch nicht als identische Verteilungen. Damit wird die Ursächlichkeit der Abweichungen von bestimmten Merkmalen von einer Zufallsverteilung und der Wandel solcher Merkmalsverteilungen zu einem der vorrangigen Untersuchungsziele der Epidemiologie. In diesem Zusammenhang sei darauf hingewiesen, dass die „Erzeugung“ von Normalverteilungen, zudem mit geringer Varianz, kein Ziel einer interventiven epidemiologischen Handlungspraxis ist. Ob dies dennoch ein Ziel sein könnte, bedarf komplexerer wissenschaftlicher Grundlagen, z.B. auch der der Sozialmedizin.

Wenn sich also die gesundheitlichen Gegebenheiten menschlicher Bevölkerungen mit dem Ablauf der Zeit ändern, bedeutet dies nichts anderes als die Bestätigung der Hypothese, dass diese Veränderungen Ursachen haben müssen, die auf die Bevölkerung wirken. Ein wichtiges Untersuchungsziel der Epidemiologie ist folglich die Ermittlung jener Ursachen, die die Differenz zwischen einer Zufallsverteilung und den empirisch vorfindlichen realen Verteilungen erklären.

Beispiel:

Da aus der historischen Perspektive die Datenlagen bezüglich der Sterbewahrscheinlichkeiten in einer Reihe von Bevölkerungen besser als die zu Erkrankungswahrscheinlichkeiten sind, sind die Sterbewahrscheinlichkeiten, also die fatalen Folgen von Krankheit, zur Erläuterung des Problems gut geeignet. Die beiden nach Gompertz und Makeham bezeichneten Terme der Sterbewahrscheinlichkeit (siehe Heft 2) beschreiben die exponentielle Zunahme der Sterbewahrscheinlichkeiten mit dem Altern mittels zweier Komponenten, nämlich einer altersunabhängigen

„Makeham Komponente“ und einer altersabhängigen „Gompertz Komponente“. Das Gompertz-Makeham Gesetz geht davon aus, dass mit dem Altern das Makeham Term sukzessive an Gewicht verliert, in den höheren Altersjahren die Sterbewahrscheinlichkeit also zunehmend und schließlich weitgehend allein vom biologischen Altern und nicht von äußeren Einflüssen abhängig ist. Dies könnte dann auch so gedeutet werden, dass einer dem Menschen wesenseigenen Sterblichkeitsfunktion (mindestens) eine zweite modifizierende Funktion gleichsam aufgesetzt ist, die die von „außen“ einwirkenden Einflüsse beschreibt. Der Altersverlauf der Wahrscheinlichkeit zu sterben, folgt also nicht allein dem Zufall. Dies gibt somit einen nachvollziehbaren Verweis auf die Beeinflussbarkeit der Lebensdauer-Verteilung in Bevölkerungen, was mit ihrem messbaren Wandel auch empirisch eindrucksvoll zu bestätigen ist. Die Analogie wird jedoch schwierig, wenn die Frage beantwortet werden soll, ob es auch bei Krankheiten, und wenn dies so ist, bei welchen, einen zufälligen und einen von der Lebensumwelt bestimmten Anteil gibt. Diese Diskussion findet unter den Stichworten „genetisch bedingt“ oder „umweltbedingt“ statt. Sie ist in ihren Konsequenzen entsprechend folgenreich, da die Interventionsansätze entweder in der „genetischen Reparatur“ oder im Wandel der Lebensbedingungen gesucht werden müssen.

Nach eigenen Untersuchungen an schwedischen Kohortendaten und deutschen Querschnittsdaten, waren die bisherigen geschichtlichen Lebensdauergewinne vor allem Folge der Minimierung des Makeham-Terms. Bis mindestens in die 1970 Jahre blieb das Gompertz-Term in den skandinavischen und in den meisten westeuropäischen Bevölkerungen weitgehend stationär. Seither ändern sich jedoch auch die Bedingungen für die Gompertzapproximation. Derzeit sinken (wenngleich geringfügig) auch die Sterbewahrscheinlichkeiten im Bereich und oberhalb der mittleren Lebensdauer. Jenseits etwa des 85. Lebensjahres ist das Gompertz-Term dann jedoch wieder stationär. Dieses Phänomen beschreibt Fries 1980 auch als „Compression of Morbidity“ (Fries, J. F. (1980): *Aging, Natural Death, and the Compression of Morbidity. New England Journal of Medicine. 303 (3): 130–5*).

Es ist empirisch überzeugend belegt, dass sich die „durchschnittliche“ Gesundheit vieler Bevölkerungen erheblich gebessert hat und weiterhin verbessert. Es lässt sich also eine langfristige tendenzielle Approximation an Zufallsverteilungen der Gesundheit an einen biologisch determinierten, wenngleich unbekannten Grenzwert vermuten.

1.1.4 Wahrscheinlichkeiten für die Richtigkeit/Falschheit einer Hypothese

Das Ziel analytischer epidemiologischer Forschung ist die Annahme oder die Ablehnung von Hypothesen.

Sofern sie nicht nur der Erkundung und der Beschreibung von Sachverhalten dient

(deskriptive Epidemiologie), prüft die analytische Epidemiologie die Übereinstimmung einer Annahme mit der Realität. Die Entscheidung hierüber erfolgt durch die Prüfung, ob eine Verteilung von einer anderen nicht zufällig unterschieden ist. Diese Entscheidung wird mittels Signifikanztests auf einem vordefinierten Signifikanzniveau getroffen. Die Feststellung von „Signifikanz“ bedeutet dann „nicht zufallsbedingt“ und ist ggf. weiter untersuchungs- und interpretationsbedürftig oder -fähig.

Für ein epidemiologisches Forschungsszenario ist folglich die Ausarbeitung eines Designs, das die Prüfung einer Studienhypothese ermöglicht grundlegend. Hierbei ist stets zu prüfen, was eine gemessene oder eine zu messende Variable tatsächlich abbildet. Im Falle der Messung von Krankheitsereignissen ist die Qualität der Diagnosestellung somit eine entscheidende Eingangsvariable. Es ist zu beachten, dass die Entscheidung, ob ein Fall von Krankheit richtig einer Diagnose zugeordnet ist, zumeist nicht durch das Forschungsteam, sondern für die Vielzahl der Fälle zumeist durch viele verschiedene behandelnde Ärzte vorgenommen wird. Es ist also immer zu prüfen, ob die Häufigkeiten von Diagnosestellungen von den „wahren“ Häufigkeiten des Krankseins abweichen. Als Standard können hier die Sensitivität, die Spezifität und ggf. der positive und der negative prädiktive Wert berechnet werden.

Die Sensitivität und die Spezifität von diagnostischen Zuordnungen von Erkrankungen zu einer Krankheitsbezeichnung sind für die Epidemiologie immer bedeutsam. Dabei hängen Sensitivität, Spezifität für prädiktive Studien auch von der Prävalenz ab. (Galen, R.S.; Gambino, S.R.: *Beyond Normality: Predictive Value and Efficiency of Medical Diagnoses*, Churchill Livingstone, 1976 ISBN 10: 0471290475 ISBN 13: 9780471290476).

1.1.5 Ungleichheit in der Epidemiologie

Das Krankwerden, dann auch als krank erkannt zu werden, angemessene Hilfe zu erfahren und ggf. unabwendbare Krankheitsfolgen bewältigen zu können, sind in Bevölkerungen ungleich verteilt.

Solche Verteilungen können durch den Zufall (Schicksal) nicht erklärt werden. Im Gedankenexperiment kann eine vollständig gleiche oder eine völlig zufällige Verteilung von Ereignissen oder Zuständen nur angenommen werden, wenn alle Menschen exakt gleich wären und stets unter identischen Bedingungen lebten. Eine solche Annahme ist realitätsfremd und steht auch im Widerspruch zu den elementaren Voraussetzungen des sozialen Lebens von homo sapiens. Unter einer solchen Bedingung wäre jede soziale Evolution auszuschließen.

So sind in Bevölkerungen auch die gesundheitlichen Gefährdungen und ihre Folgen stets ungleich verteilt. Sie sind in einer historischen Perspektive zugleich Ursache und Folge sozial evolutiver Vorgänge. Solche Ungleichheiten sind also nicht die Ausnahme oder etwas Besonderes. Sie sind auch nicht immer Zeichen eines

Systemversagens. Ungleichverteilungen sind die direkten und indirekten Folgen der aktiven gesellschaftlichen wie produktiven Austauschverhältnisse der Menschen, ihrer individuellen wie sozialen Entwicklung und letztlich auch eine Schlüsseldeterminante des historischen Wandels der gesundheitlichen Eigenschaften von Bevölkerungen. Sozialer Wandel - gleich wie dieser hinsichtlich seiner Folgen jeweils zu bewerten sei - ist unvermeidlich mit der Differenzierung der gesundheitlichen Verhältnisse verbunden. Entsprechende ungleiche Verteilungen und die Wandlungsfähigkeit bedingen einander.

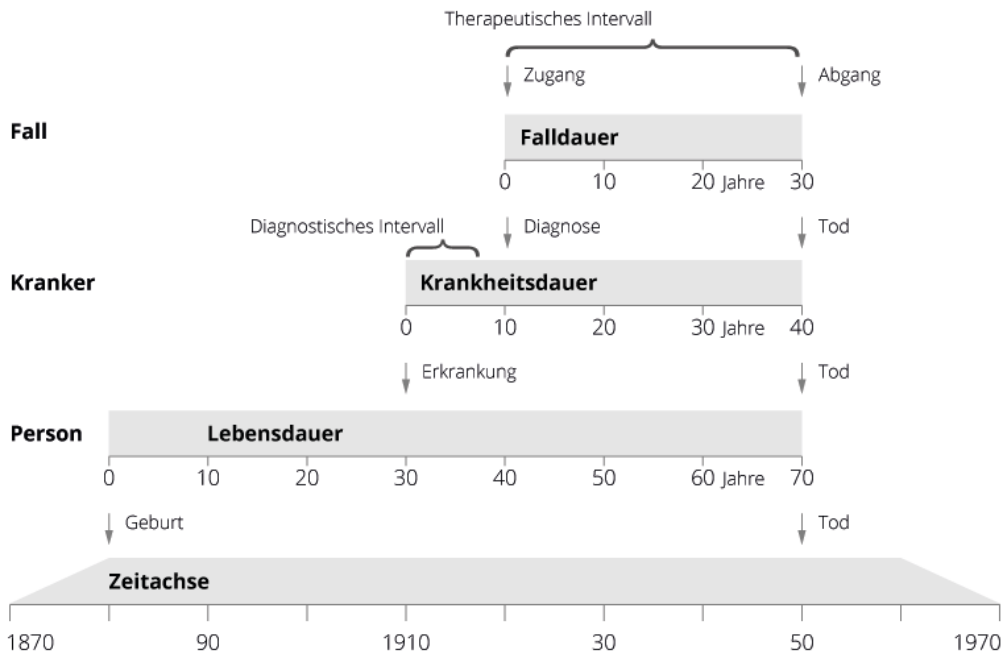


Abb. 2: Beziehung von Falldauer, Alter und epochaler Zeit

Aus diesem Grunde ist das Verständnis der qualitativen Richtung und der Ursachen des Wandels gesundheitlicher Probleme eine wichtige Voraussetzung für die Bewertung solcher Ungleichverteilungen. Diese beziehen sich immer auf die Zeit, nämlich

1. auf die zeitlichen Parameter pathogenetischer Prozesse
2. auf das Lebensalter bzw. die bereits erreichte Lebensdauer und
3. auf die epochale Zeit

1.1.6 Epidemiologische und klinische Studien

Epidemiologie ist die Lehre von den Epidemien und ihrem Gegenteil, den Regressionen. Klinische Studien erzeugen Wissen über die Diagnostik und die Behandlung von individuellen Patienten.

Epidemiologische und klinische Studien folgen unterschiedlichen Forschungszielen. Gemeinsam ist ihnen die Nutzung von mathematischen Methoden.

In epidemiologischen Studien ist die Ermittlung der realen Heterogenität der gesundheitlichen Eigenschaften von Bevölkerungen das typische Studienziel. Die untersuchte Bevölkerung muss folglich der realen Bevölkerung entsprechen, diese also repräsentieren. Repräsentativität im Verhältnis zur untersuchten Bevölkerung ist deshalb eine Schlüsselbedingung epidemiologischer Studien. Die Repräsentativität der Studiengruppen ist dann das wichtigste Qualitätsmerkmal epidemiologischer Studien.

Klinische Studien untersuchen die Effekte von Intervention in gegenüber der Bevölkerung selektierten Studiensamples. Die untersuchten Individuen sind möglichst homogen und unterscheiden sich idealerweise nur bezüglich der untersuchten Variablen. Um systematische Verzerrungen zwischen der Interventionsgruppe und der Kontrollgruppe zu vermeiden, werden die Teilnehmer den Gruppen zufällig zugeordnet (Randomisierung). Die Randomisierung der Interventions- und der Kontrollgruppe ist das wichtigste Qualitätsmerkmal klinischer Studien.

Epidemiologische und klinische Studienpopulationen unterscheiden sich also nicht notwendig durch ihre Umfänge, sondern durch ihre Struktur, bzw. die Dynamik ihrer Strukturen.

Repräsentativität und Randomisierung definieren die entscheidenden Unterschiede von epidemiologischen und klinischen Studien.

Für beide Ansätze, also die Repräsentativität und die Randomisierung existieren dann auch jeweils spezifische Auswahltechniken und nachfolgend unterschiedliche Interpretationsbedingungen der Studienergebnisse.

Beispiel:

Ob ein neu entwickeltes Medikament wirksam ist oder gegenüber einem eingeführten Medikament einen zusätzlichen klinischen Nutzen hat, ist eine typische klinische Fragestellung. Wie viele Menschen einer realen Bevölkerung ggf. von diesem

Medikament profitieren würden, ist hingegen eine epidemiologische Fragestellung, die durch eine klinische Studie nicht entschieden werden kann.

Im populären Verständnis wird oft angenommen, die Dokumentation und klassifizierende Zuordnung von Sachverhalten, also das Beschreiben, Messen und Zählen von Ereignissen und Zuständen sei die Aufgabe von Statistik. Tatsächlich ist die Statistik jedoch ein Teilgebiet der Mathematik, die u.a. auch für epidemiologische wie für klinische Studien genutzt wird. Die Mengenlehre und die Wahrscheinlichkeitstheorie sind dabei jeweils wichtige Grundlagen des wissenschaftlichen Arbeitens. Die Gemeinsamkeit der Nutzung der mathematischen Statistik kann jedoch nicht die Annahme begründen, beide Studientypen seien dem Ziel und der Arbeitsweise nach wesensgleiche wissenschaftliche Arbeitsfelder.

Das bedeutet auch, dass Studierende der Medizin und dann später als Ärztinnen und Ärzte arbeitende Entscheider und Gutachter eines hinreichenden statistischen Grundverständnisses bedürfen, um quantitative Forschungsergebnisse verstehen, kommentieren und praktisch nutzen zu können. Ist das nicht der Fall, bleibt ihnen nur, anderen Experten zu vertrauen. Vertrauen ist aber immer auch die Preisgabe einer eigenen und durch Wissenschaft begründeten (evidenzbasierten) Entscheidungskompetenz.

Die wichtigsten Unterschiede zwischen epidemiologischen und klinischen Studien sind:

Die Epidemiologie fragt nach den gesundheitlichen Eigenschaften von Bevölkerungen. Auf diese mit den Systemen und Infrastrukturen der Prävention und der Krankenversorgung reagieren zu können, ist das praktische Ziel. Es wird also in der Regel darum gehen, die tatsächliche Heterogenität der untersuchten Bevölkerung abzubilden und Wirkungen unter den Bedingungen dieser Heterogenität zu ermitteln. Folglich sind die Repräsentativität der Studienbevölkerung im Verhältnis zur Zielbevölkerung die entscheidende Qualitätsfaktoren in der Epidemiologie. Aus inhaltlicher Sicht steht also die bevölkerungsbezogene Relevanz eines Ergebnisses für die Zielbevölkerung im Vordergrund epidemiologischer Forschung.

Ihre Forschungsgegenstände sind z.B. Gesundheitsgefahren und -chancen, Epidemien und Regressionen, Hilfebedarfe und medizinische, physische, psychische oder soziale Krankheitsfolgen.

Eine Bevölkerung ist in der Perspektive der Epidemiologie eine Gesamtheit von Menschen, die durch ihren Umfang und ihre Struktur sowie ihre Lebensweise gekennzeichnet ist. Diese Struktur können soziale Merkmale, biologische Merkmale oder auch klinische Merkmale je nach Studienziel beschreiben. Die einzelnen Strukturmerkmalen zurechenbaren Personen sind dann jeweils Teilbevölkerungen.

Epidemiologische Studienbevölkerungen können vollständig (Totalerfassungen) zu meist auf der Basis von gesetzlich geregelten Datenerfassungssystemen) sein oder durch Stichproben gebildet werden. Stichproben sind Abbilder einer Zielbevölkerung

(Grundgesamtheit). Sie entstehen durch die Zufallsauswahl von Elementen der Grundgesamtheit. Kenntnisse über die Grundgesamtheit sind somit Voraussetzungen für Stichprobenziehungen.

Der Wert stichprobenbasierter Aussagen hängt davon ab, wie gut diese die Grundgesamtheit abbilden. Repräsentativität bedeutet dann, dass die Stichprobe mit einem abschätzbaren Fehler die reale Bevölkerung widerspiegelt. Da die Repräsentanz einer Stichprobe von der Fragestellung einer Studie abhängt, kann eine repräsentative Stichprobe ohne Bezug zur Fragestellung nicht gezogen werden. Die Ergebnisse stichprobenbasierter Studien sind dann innerhalb der jeweiligen Konfidenzintervalle gültig. Stichprobenaussagen dürfen nicht mit dem Begriff der Hochrechnung verwechselt werden. Sie sind auch nicht weniger "richtig" als Totalerfassungen. Sie sind immer in den Grenzen der Konfidenz exakt richtig. Ohne Angabe der Fragestellung und ohne Angabe der Konfidenz kann es deshalb keine korrekten stichprobenbasierten Angaben über Bevölkerungen geben.

Aktuelle Entwicklungen zeigen die wachsende Bereitschaft eines Teils der Bevölkerungen, gesundheitsrelevante Daten dem Zugriff professioneller „data broker“ verfügbar zu machen. Das wirft u.a. die Frage nach dem Wert solcher Daten auch für wissenschaftliche Fragestellungen auf. Die Antwort hierauf verlangt die Klärung des jeweiligen Kontextes. Für epidemiologische Fragestellungen wäre jeweils zu klären, ob solche Daten die Grundgesamtheit einer Bevölkerung repräsentieren. Eine solche Klärung ist praktisch nie möglich, da zumeist nie offenzulegen ist, wie Datenhändler in den Besitz von Daten gekommen sind. Das solche Entwicklungen rechtfertigenden Theorem besagt, Bevölkerungen seien mittels der Konzepte der Schwarmintelligenz hinreichend gut zu beschreiben, es bedürfe folglich nicht der Kenntnis der biologischen und sozialen Struktureigenschaften von Bevölkerungen. Damit würden sich dann auch Anforderungen an die Repräsentativität von Studienbevölkerungen erübrigen.

In der klinischen Forschung steht das Kranksein des individuellen Patienten im Mittelpunkt des Interesses. Ziel ist es, wirksame (standardisierte) Interventionsstrategien zu entwickeln und zu prüfen. Es geht in der Regel also darum, durch geeignete Auswahltechniken ein Höchstmaß an Homogenität der Studienpopulation zu erzeugen. Die erfolgreiche Randomisierung der Studienbevölkerung ist also der entscheidende Qualitätsanspruch. Aus methodischer Sicht geht es darum, zwischen Interventions- und Kontrollgruppe Verteilungsunterschiede bezüglich der Zielkriterien (Outcomes) zu ermitteln. Die Studienergebnisse sind dann in dem Maße zu verallgemeinern, wie sie sich auf Personen beziehen lassen, die den Studiengruppen gleichen oder ähneln.

Die Randomisierung und nicht die Repräsentativität ist also das maßgebliche Qualitätskriterium klinischer Studien. Die Randomisierung umfasst Verfahren der zufälligen Zuordnung von Personen jeweils zu den zu vergleichenden Studiengruppen. In klinischen Studien „stört“ die Heterogenität der Probanden. Wesentliche

Heterogenitätsmerkmale von Bevölkerungen können das Ergebnis einer klinischen Studie also verzerren, bspw. durch das sog. Confounding.

Wenn in epidemiologischen Studien Mengen von als gleich bezeichneten gesundheitlichen Zuständen Ausgangspunkt einer Analyse sind, muss immer auch gefragt werden, ob hier tatsächlich gleiche Zustände zusammengefasst sind. Das ist oft nicht zu gewährleisten. Das Ausmaß dieser „Unschärfe“ muss in epidemiologischen Studien ermittelt und bei der Interpretation der Ergebnisse berücksichtigt werden.

Der praktische Nutzen epidemiologischer Studien liegt z.B.

1. in der Definition von Zielgruppen von Präventionsprogrammen (Identifizierung von epidemiologischen Potenzialen oder Risikogruppen)
2. in der Ermittlung von Versorgungsbedarfen (Bedarfs- und Allokationsstudien)
3. in der Effektschätzung von (z.B. präventiven) Interventionen in einer Bevölkerung (Interventionsstudien)
4. in den Kalkulationsbedürfnissen von Äquivalenzversicherungen (Risikoklassifikationen).

Der praktische Nutzen klinischer Studien liegt hingegen z.B.

1. in der Prüfung und Bewertung therapeutischer Konzepte
2. in der Risikobewertung von Therapien
3. in der Identifizierung der jeweiligen Patientengruppen, die von einer Therapie profitieren oder nicht profitieren (intention to treat analysis, attributable therapeutic fraction analysis).

2 Grundbegriffe der Epidemiologie

2.1 Epidemien und Regressionen

Jede Zunahme von Zustandsmengen (Prävalenz) ist eine Epidemie. Jede Abnahme von Zustandsmengen ist eine Regression.

Die Epidemiologie untersucht Mengenänderungen von Fällen und Personen mit gesundheitsrelevanten Eigenschaften. Dies sind definierte Gefährdungen (Expositionen), Krankheiten und Krankheitsfolgen. Bei nicht wiederholbaren Krankheiten entspricht die Fallmenge der Personenmenge. Sofern sich die jeweiligen Mengen auf Sekundärdaten beziehen, sind die dokumentierten Fallmengen (derzeit) praktisch nie zur Person zusammenzuführen. Es bleibt dann unklar, ob eine Epidemie/Regression sich auf Fallmengen oder auf die tatsächlich betroffene Bevölkerung bezieht.

Die von den Autoren des Heftes angebotene Definition schließt eine Reihe von Festlegungen ein:

1. Epidemien und Regressionen sind Bewegungen von Zustandsmengen (Prävalenz). Sie können folglich für das Gleichgewicht von Zugängen und Abgängen zu den Bestandsmengen nicht definiert werden.
2. Das Gleichgewicht ist unabhängig vom Umfang der Prävalenz.
3. Epidemien und Regressionen können grundsätzlich auf alle Mengenänderungen von gesundheitsrelevanten Zuständen, also auch für gefährdete Personen, für Kranke und für Personen, die an Krankheitsfolgen leiden, beschrieben werden. Das gilt analog für Fallmengen.
4. Epidemien und Regressionen einzelner gesundheitsrelevanter Zustände, z.B. übertragbare und nicht übertragbare Gefährdungen und Krankheiten, sind jeweils als deren Sonderfälle aufzufassen.
5. Epidemien und Regressionen beschreiben die jeweilige Richtung der Zustandsänderungen der Prävalenzen.
6. Gleichgewichtsbedingungen sind gegeben, wenn die Zugänge zu den Zustandsmengen und die Abgänge aus ihnen gleiche Beträge haben.

Die Phänomene "Epidemie" und "Regression" werden nicht gleichlautend definiert. Während einige Autoren eine Epidemie als *"das Auftreten einer bestimmten Krankheit gleichzeitig bei einer großen Anzahl von Menschen"* (<https://dictionary.cambridge.org/Zugriff> 29.01.2022) definieren, bezeichnen andere, so wie auch wir, eine Epidemie exakter als die Zunahme der Prävalenz.

Es besteht aus unserer Sicht allerdings ein Interesse an einem besseren Verständnis dessen, was Epidemien und Regressionen sind. Einer der Gründe ist die Notwendigkeit, Personen und epidemiologische Fälle zu unterscheiden, da eine Person (Träger) mehrerer „Fälle“ sowohl gleicher wie auch unterschiedlicher „Fälle“ sein kann.

Die Autoren der *Principles of Epidemiology* (3. Auflage, 2012, Atlanta, Georgia; herausgegeben vom Centers for Disease Control and Prevention) beschreiben Epidemien als *"... das Auftreten von mehr Krankheitsfällen, Verletzungen oder anderen Gesundheitszuständen als erwartet in einem bestimmten Gebiet oder unter einer bestimmten Personengruppe während eines bestimmten Zeitraums" ... „In der Regel wird davon ausgegangen, dass die Fälle eine gemeinsame Ursache haben oder in irgendeiner Weise miteinander zusammenhängen."* (Übersetzung durch die Autoren.) Im Ergebnis einer näheren Textanalyse kommen wir zu dem Schluss, dass die Begriffswahl „Auftreten“ die Menge gleichartiger „Zustände“, also die Prävalenz meint.

Es ist ein Dilemma, dass epidemiologische Fälle mit wenigen Ausnahmen nur dann zu diagnostizieren und zählen sind, wenn sie bereits existieren. Es gibt nur wenige Beispiele, dass alle vorkommenden „Fälle“ in dem Moment, in dem sie entstehen diagnostiziert werden oder diesem Zeitpunkt zumindest nahe diagnostiziert werden. Das „diagnostische Intervall“ zwischen dem Zeitpunkt des „Entstehens“ eines neuen Falls und dem der „richtigen“ Diagnosestellung ist für die Identifizierung von Epidemien, aber auch von Regressionen von grundlegender Bedeutung. Dies gilt insbesondere dann, wenn sich die Verteilung aller fallspezifischen Diagnoseintervalle im Laufe der Zeit systematisch ändert. Ist dies der Fall, kann dies erhebliche quantitative Auswirkungen haben, die dann die Interpretation der Bewegung der Fallmengen verzerren. Es ist also zwischen dem Zeitpunkt zu dem Menschen erkranken, und dem Zeitpunkt, zu dem sie diagnostiziert werden, zu unterscheiden. Bezogen auf eine Bevölkerung sind diese Intervalle über ihre zeitliche Verteilung zu beschreiben. (Niehoff, J.-U. *Sozialmedizin systematisch*. 3. Auflage UNI-MED Verlag, Bremen, Boston, London, 2011, S. 65, ISBN 978-3-8374-1283-3).

Epidemiologen diagnostiziert in der Regel keine einzelnen Fälle einer Krankheit und behandeln sie auch nicht. Sie sind auf Daten angewiesen, die nur unter sehr speziellen Studiendesigns auch standardisiert erhoben werden. Die Bewertung nicht nur der diagnostischen Methoden, sondern auch die ihrer Anwendung und Zugänglichkeit sind für die epidemiologische Analytik folglich von zentraler Bedeutung.

Jede Bewegung der Prävalenz kann nur zwei Richtungen aufweisen, nämlich ihre "Zunahme" oder ihre "Abnahme". Und diese Dynamik kann nur als eine Bewegung zum „Schlechteren“ oder zum „Besseren“ beurteilt werden. Solche Urteile gelten stets nur für konkrete Bevölkerungen und Teilbevölkerungen, können für diese in Richtung und Geschwindigkeit also jeweils unterschiedlich sein. Ist dies der Fall, sind Ungleichheiten, ggf. Benachteiligungen die Folge.

Die Bewegungen „Epidemie“ und „Regression“ sind nicht mit den Bewegungen der jeweiligen Inzidenzen, bzw. mit denen der quantitativen Bewegung von „Ereignissen“ gleichzusetzen. Es ist also bei der Identifizierung von Epidemien/Regressionen zwischen den Bewegungen von Ereignissen und Zuständen strikt zu unterscheiden. Aus der Veränderung einer Prävalenz kann nicht auf eine Veränderung von krankheitsspezifischen Gefährdungen geschlossen werden. Die spezifischen Bedeutungen werden durch die Tatsache definiert, dass nur eine bestimmte Gruppe von Krankheiten eine

kurze und konstante Falldauer aufweist und dann Inzidenzen und Prävalenzen gegenseitig als Schätzer dienen können.

Die Erfolge der Medizin in Bevölkerungen mit (schon) hoher mittlerer Lebensdauer gehen zu einem großen Anteil auf Effekte zurück, die der Beeinflussung der diagnostischen und der therapeutischen Intervalle folgen. Die Fortschritte in der Medizin ermöglichen zudem einem wachsenden Anteil von Menschen ein langes Leben, auch wenn sie an Krankheiten leiden, die noch nicht zu heilen sind. Gleichzeitig ist der bessere Zugang zur medizinischen Versorgung ein wesentliches Merkmal der Evolution von Systemen der Krankenversorgung. Beide Fortschritte können die Bewegung der dokumentierten epidemiologischen Fallzahlen in der Weise beeinflussen, dass sie "falsche" oder nur „scheinbare“ Epidemien und Regressionen auslösen.

Für die Beschreibung von Epidemien und Regressionen gilt

$$P \sim I \times D$$

Im besonderen Fall des Gleichgewichts der Inzidenz mit den Abgängen oder dann, wenn die Dauer (D) einer Krankheit ihr konstantes Merkmal wäre, ist diese Gleichung

$$P = I \times D$$

P = Prävalenz

I = Inzidenz

D = Dauer

Die Mengen der prävalenten Fälle sind als im Gleichgewichtszustand befindlich zu beschreiben, wenn sich weder die Inzidenz noch die mittlere Dauer der spezifischen Krankheit ändert. Unter dieser Bedingung sind die Mengen der Zugänge zur und der Abgänge aus der Prävalenz gleich.

Pandemien sind von Epidemien allein durch ihre Globalisierung unterschieden. Sie sind also Aggregationen von Epidemien, so wie Epidemien Aggregationen von Mikroepidemien sind.

Beispiel:

Die Sars-CoV-2-Pandemie im Jahr 2020 und den folgenden Zeiten war eine Aggregation vieler regionaler Epidemien. Das kann nicht darüber hinwegtäuschen, dass die nationalen Epidemien der USA, oder Indiens, der VR China oder der Philippinen usw. in ihren spezifischen Ursächlichkeiten, Abläufen und Folgen identisch wären. Dies gilt auch für die vielen Mikro-Epidemien innerhalb dieser Länder. Die Anwendung epidemiologischer Erkenntnisse auf die nationalen Gegebenheiten erfordert

bei der Bekämpfung und Prävention einer nationalen Epidemie die Kenntnis der jeweils konkreten Bedingungen in den jeweiligen Staaten.

2.2 Epidemiologische Potenziale

Epidemiologische Potenziale bezeichnen die Teile von Bevölkerungen, die sich nach dem Grad der Gefährdung durch eine Exposition und deren Folgen voneinander unterscheiden, bzw. diskriminieren lassen.

Soweit wir wissen, wurde der Begriff "epidemiologisches Potenzial" erstmals 1987 bei der Analyse der AIDS-Epidemie verwendet und diskutiert (Niehoff, J.-U., Sönnichsen, N. (1987): AIDS - Merkmale einer echten Epidemie. Ztschr. klin. Med. 42:2141-2145).

Die inhaltliche Bedeutung folgt aus der Klassifikation der Gefährdeten, der Betroffenen und der an Krankheitsfolgen Leidenden. Ihre methodische Bedeutung liegt in der Funktion jeweils die sachlich logischen Denominatoren für die Kalkulation von Anteilsziffern und von Wahrscheinlichkeiten zu sein.

Beispiel:

Soll ermittelt werden, wie hoch der Anteil der Raucher an einer Bevölkerung oder einer ihrer Teilbevölkerung ist, ist das entsprechende epidemiologische Potenzial jeweils diese Bevölkerung. Soll hingegen ermittelt werden, wie groß die Wahrscheinlichkeit ist, dass eine Teilbevölkerung von Nichtraucherern mit dem Tabakkonsum beginnt (oder analog ihn beendet) müssen die jeweiligen Ausgangsmengen der Nichtraucher, bzw. der Tabakkonsumenten ermittelt werden. Das epidemiologische Potenzial der Tabakkonsumenten ist dann das Potenzial, auf das sich Wahrscheinlichkeitskalküle für Folgekrankheiten logisch richtig beziehen können. Das Beispiel ist vollständig auch auf Infektionskrankheiten zu übertragen. Auch hier ist einmal von Interesse zu wissen, wie hoch der Anteil infizierter Personen an der Gesamtbevölkerung oder jeweils an ihren strukturellen Gliederungen ist. Es ist dann logisch korrekt, die im Verlauf einer Epidemie dokumentierten neuen Fälle jeweils nur auf das Potenzial der nicht Immunisierten zu beziehen. Ebenso kann es von erheblichem Interesse sein, die Wahrscheinlichkeit zu ermitteln, mit der Personen unterschiedlicher Immunität erkranken, genesen oder versterben.

Solche Teilbevölkerungen, z.B. charakterisiert durch gleiche Gefährdungen, sind dann als "epidemiologisches Potenzial" zu bezeichnen. Die grundsätzliche Bedeutung solcher Potentiale ergibt sich aus der Inhomogenität von Bevölkerungen in Bezug auf deren biologische und soziale Merkmale. Dazu gehört die Ungleichheit von Teilen einer Bevölkerung, mit Expositionen in Berührung zu kommen, dann auch zu erkranken und, wenn sie erkranken, schwerwiegende Folgen zu erleiden oder auch nicht.

Epidemiologen benötigen ein präzises Bild von der Unterschiedlichkeit der gefährdeten Untergruppen einer Bevölkerung. Nicht nur die Quantitäten, sondern auch die sozioökonomischen Strukturmerkmale der verschiedenen epidemiologischen Potenziale sind von Bedeutung. Die relevanten Strukturmerkmale sind z.B. Geschlecht, Alter, Verhalten, Lebensstil, Komorbiditäten, Bildung, Beruf, soziale Ressourcen und alle anderen Merkmale, die die vertikale und horizontale soziale Varianz dieser verschiedenen Potenziale beschreiben. Analoge Begriffe sind Cluster oder auch Risikogruppen.

Es gibt keine Epidemiologie ohne eine Klassifizierung der Populationen und ihrer Untergruppen.

Die Gruppierung von Elementen als gleich bedeutet nicht, dass sie identisch sind. Jedes einzelne Element kann durch viele verschiedene Konzepte und für viele Zwecke klassifiziert werden. Dies gilt auch für die Klassifizierung von Epidemien und Regressionen, für die Klassifizierung ihrer spezifischen epidemiologischen Potenziale, für die Klassifizierung der gefährdeten Personen oder für die Klassifizierung der Folgen von Epidemien und Regressionen, usw. Das Hauptziel besteht darin, sowohl die Subjekte (wer) als auch die Objekte (was) zu klassifizieren.

Es ist eine Voraussetzung für die Klassifizierung von Elementen, dass sie nach klar definierten Kriterien gruppiert werden. Problematisch ist allerdings die Klassifizierung von Personen nach ihrem Alter. Die medizinische Wissenschaft verfügt über kein praktikables Maß für das biologische Alter von Personen. Sie verwendet zu dessen Schätzung das kalendarische Alter. Das kalendarische Alter ist jedoch kein guter Schätzer für das biologische Alter. Dies ist ein Schwachpunkt der Epidemiologie, aber viele medizinische Studien teilen dieses Dilemma.

Für den Begriff des epidemiologischen Potenzials hat sich auch der Begriff „Risikogruppe“ etabliert. Entsprechend sind die Merkmale, die solche Gruppen unterscheiden Risikomerkmale, Risikoindikatoren oder Risikofaktoren. Da der Begriff „Risiko“ immer auch eine von dem Begriff „Chance“ zu unterscheidende Wertung beinhaltet, gehen die Autoren dieses Heftes mit dem Begriff „Risiko“ zurückhaltend um. Nach unserer Auffassung sind analytische und Sachverhaltsermittlungen wie sie die Epidemiologie zu verantworten hat, von Werturteilen zu trennen.

Personen mit gleichen oder ähnlichen Gefährdungsprofilen sind „Risikogruppen“. Die den Gefährdungen des „Gesundseins“ assoziierten Variablen werden als „Risikofaktoren“ bzw. Risikoindikatoren bezeichnet.

Entsprechende Merkmale liegen in Form von absoluten Häufigkeiten oder als sog. Anteilsziffern und als Häufigkeitsunterschiede (odds ratios), eher selten als

Wahrscheinlichkeitskalküle vor. Solche „odds ratios“ sind keine spezifischen Merkmale eines Potenzials und für Vergleichszwecke nicht geeignet. Absolute Häufigkeiten erfüllen ihren Zweck, wenn sie für Bedarfs- und Allokationszwecke genutzt werden sollen. Für räumliche und zeitliche Vergleiche sind sie ungeeignet. Hierzu bedarf es mindestens der Relativierung von absoluten Häufigkeiten durch den Bezug auf ihre jeweiligen epidemiologischen Potenziale und für Vergleichszwecke zudem der Standardisierung der Vergleichsbedingungen (siehe hierzu das Kapitel 4 des Heftes).

Epidemiologische Studien zur Verteilung und zum Verteilungswandel gesundheitsrelevanter Eigenschaften von Bevölkerungen zielen letztlich auf die Charakterisierungen von Gefährdungen für die Gesundheit bzw. von Chancen, diese zu bewahren. In der Praxis, zumal in der medizinischen Praxis individueller Beratung und therapeutischer Entscheidungen, sind die dem Risikobegriff unterlegten statistischen Konstruktionen nur bedingt nutzbar, weil Vergleiche dem Einzelfall nur bedingt helfen. Zumal Wahrscheinlichkeiten für den Einzelfall nicht definiert werden können und Risikoausagen nicht von schwierigen Entscheidungen über ihre Akzeptanz im klinischen Einzelfall befreien. Es geht dabei eher um die qualitative Bewertung und Abwägung konfligierender Möglichkeiten. Das Argument „Möglichkeit“ hat in der Praxis zwei Aspekte:

1. Das Argument bezieht sich auf die allgemeine qualitative Erfahrung, dass eine Wirkung eintreten kann.

Beispiel:

Unabhängig von Alter und Geschlecht können alle Menschen einen Schlaganfall erleiden. Die Häufigkeit solcher Ereignisse variiert in Abhängigkeit von zusätzlichen Merkmalen, deren Träger aus praktischen Erwägungen klassifiziert werden können (Risikogruppen).

2. Es greift auf konkrete quantitative Schätzungen der statistischen Wahrscheinlichkeit für das Auftreten eines Ereignisses zurück.

Beispiel:

Von 100 Personen, die täglich 20 Zigaretten rauchen und damit im Alter von 20 Jahren begonnen haben, können bis zum Alter von 70 Jahren, also nach maximal gemeinsamen 1.825.000 Expositionstagen, 10 an einem Lungenkarzinom verstorben sein.

In einer individuellen Beratung ist eine solcher Feststellung, empirische Richtigkeit vorausgesetzt, nicht sinnvoll zu argumentieren. Es bedürfte mindestens des Argumentes, dass Nichtraucher seltener erkranken oder länger leben. Beides sind aber

dann Feststellungen, die nicht unabhängig von den Gefährdungen der Nichtraucher in der Gesamtbevölkerung Überzeugungskraft haben.

Risikoaussagen sind also immer vergleichende Aussagen und damit frei wählbar, für den Ratgeber wie für den Beratenen. Die Risikoakzeptanz ist in solchen Beratungssituationen u.U. wichtiger als die Höhe eines Risikos. Ursachen von Erkrankungen, die prospektiv erneut oder auch für andere Personen wirksam werden können, sind immer Risiken, nur ohne Relationen sind sie jedoch argumentationsschwach. Risiken sind mögliche Ursachen für eine Erkrankung. Sie können im individuellen Einzelfall wirksam werden oder auch nicht.

Das epidemiologische Argument ist also immer ein quantitatives Argument. Seine Stärken liegen eben darin, allerdings unter der Voraussetzung, die Risikowarnung richtet sich an die bezüglichen epidemiologischen Potenziale. Die Schaffung solcher auf konkrete Settings bezogenen Präventionsargumente ist die Aufgabe der interventionellen Epidemiologie. Ihre analytische Aufgabe ist die relationale Quantifizierung von Krankheitsereignissen konkreter epidemiologischer Potenziale in Bevölkerungen.

Der Begriff „Risiko“ bezeichnet somit allgemein die Möglichkeit des Eintretens eines Schadens. Der Begriff „Risikofaktor“ bezeichnet dann die Einwirkung („Exposition“), die diesen Schaden verursachen kann. „Risikofaktoren“ können sowohl Strukturmerkmale einer Bevölkerung, die deren Risikoinhomogenität beschreiben (Risikoindikatoren), als auch auf Menschen gleichsam von außen einwirkende Expositionen sein. Da Bevölkerungen hinsichtlich der Gefährdungen für die Gesundheit nicht homogen sind, ist es häufig wichtig, Gefährdungsdifferenzen zwischen einzelnen Teilen einer Bevölkerung zu identifizieren und zu quantifizieren. Aufgrund der Vielzahl der Einwirkungen auf jeden Menschen in einer Bevölkerung ist die Bewertung solcher Differenzen meist schwierig. Interaktionen von Menschen mit der sie umgebenden Welt können niemals vollständig risikolos sein. Die Bewertung von Risiken ist folglich ein Schlüsselproblem, das von der Epidemiologie allein nicht geleistet werden kann, muss sich aber immer auf relationale Aussagen berufen können. Relationale oder relativer Risiken sind also „Distanzmaße“. Sie sind so lange verschieden, wie sich Gruppen von Menschen unterscheiden.

In der Praxis des Risikomanagements und in risikoäquivalenten Krankenversicherungen steht das Interesse an Schätzungen der Anzahl neu auftretender Fälle im Vordergrund. Folglich wird eher die Frage, wie viele gegenüber einem Risikofaktor exponierte Personen krank werden bearbeitet, als die Frage, warum, ggf. sehr viele gegenüber dem Risikofaktor exponierte Personen nicht erkranken. Auch die Frage nach den Ursachen von Gefährdungen ist dann oft von nachgeordneter Bedeutung. Risiko und Chance werden also in der Realität häufig nicht symmetrisch wahrgenommen.

Risikoschätzungen dienen unterschiedlichen Interessen:

- In der Prävention bezeichnet das Risiko die Wahrscheinlichkeit einer künftigen Erkrankung, also eines Schadens für den Betroffenen.

- In der Krankenversicherung bezeichnet das Risiko die Wahrscheinlichkeit der Inanspruchnahme einer Versicherungsleistung, also eines Schadens für die Versicherung.
- In der medizinischen Versorgung bezeichnet das Risiko die Wahrscheinlichkeit einer Nebenwirkung bzw. eines nicht gewünschten Behandlungsergebnisses, ggf. also eines Haftungsfalles.

In der Entscheidungspraxis sind solche Kalküle sinnvoll, wenn ihnen ein Bewertungskonzept zugeordnet werden kann. Gleiche Sachverhalte können je nach Kontext und Interessenlage jedoch häufig sowohl als Risiko wie auch als Chance gedeutet werden.

Die vorausschauende Folge- und Schadensabschätzung (prospective modelling) ist in der Prävention, bei Therapieentscheidungen, prospektiven Finanzierungsmodellen für medizinische Leistungen oder bei den Prämienkalkulationen von for-profit Versicherungen eine regelmäßige Praxis.

Hierzu werden Prädiktoren (Wahrscheinlichkeitskalküle) ungewollter Ereignisse (z.B. Auslösung von Krankheiten) genutzt. Die Methoden sind auch bei unterschiedlichen Interessen ähnlich. Allerdings muss die Besonderheit der jeweiligen Interessen bei der Bewertung der resultierenden „Risiken“ immer mitbeachtet werden.

Folgende Fragen sind dann zu beantworten:

1. Welchem Zweck dient eine prospektive Folgeabschätzung?
2. Welche Intensität hat eine Einwirkung (force of morbidity)?
3. Wie ist ein Risiko in einer Bevölkerung verteilt?
4. Wie variant ist in einer Bevölkerung die Vulnerabilität gegenüber einer Gefährdung verteilt?
5. Können einer potenziell unerwünschten Folge (Risiko) auch gewünschte Folgen (Chancen) zugerechnet werden und jeweils von wem?

Mit der Identifizierung von Risiken entsteht die Frage, wer von diesem Merkmal betroffen, für wen dieses Merkmal also bedeutsam ist (Risikogruppe) und wer es ggf. zu verantworten hat.

Die Begriffe Risikofaktor und Risikogruppe haben einen vielfältigen Gebrauch gefunden. Hierfür gibt es vor allem zwei Gründe. Der 1. folgt den Bedürfnissen von Lebensversicherungen und Krankenversicherungen, die nach dem Prinzip der Risikoäquivalenz arbeiten und ihre Leistungen gegen die Beiträge aufrechnen. Der Beitrag wird dann von dem kalkulierten Schaden infolge der Eigenschaften der Versicherten abhängig gemacht (Risiko=Schuld). Der 2. Grund folgt den Bedürfnissen einer prädiktiven, also auf vorhersehende Handlungsorientierungen ausgerichteten interventionellen Präventivmedizin sowie präventiven öffentlich-rechtlichen Handlungsorientie-

rungen (Public Health). Beide Bedürfnisse haben die Denk- und Arbeitsweise der Epidemiologie besonders seit den frühen 1950er Jahren nachhaltig geprägt.

Risiken können logisch auf zweierlei Weise begründet werden.

1. Risiken beschreiben die Gefährdungen von Personengruppen.

Beispiel:

Von 1.000 Neugeborenen in der Bevölkerung X erkranken innerhalb der nächsten 40 Jahre 10 an einer Tuberkulose, d.h. unter den gegebenen Bedingungen beträgt die kumulierte Inzidenz, bzw. das über diese 40 Jahre verteilte Risiko 0,001.

2. Risiken beschreiben die Gefährlichkeit einer Einwirkung in Abhängigkeit von der Zeit.

Beispiel:

Je 1.000 Expositionsstunden gegenüber der chemischen Substanz Y sind 4 Erkrankungen an der Krankheit A zu beobachten; oder, bei der gegebenen Expositionstärke führen 250 Expositionsstunden durchschnittlich zu einer Erkrankung.

Ein Risikomerkm al bezeichnet oder indiziert ein Merkmal, mit dem ein gesundheitliches Risiko assoziiert ist.

Ein Risikomerkm al klassifiziert also eine Inzidenzdifferenz zwischen zwei ansonsten strukturgleichen Potenzialen als disjunkt. Hierfür ist unerheblich, ob das Ausmaß dieses Unterschiedes diesem Merkmal ursächlich zuzurechnen ist. Es ist zu beachten, dass die Strukturgleichheit den Interpretationsrahmen festlegt.

Risikofaktoren erklären ursächlich die Differenz der Inzidenz zwischen mindestens zwei ansonsten strukturgleichen Bevölkerungen.

Risikofaktoren unterscheiden sich von Risikomerkm alen dadurch, dass hinreichende Belege dafür vorliegen, dass die Inzidenzdifferenz zwischen mindestens zwei Personengruppen, von denen nur eine diesen Faktor trägt, ursächlich auf diesen zurückzuführen ist.

Beispiel:

Raucher erkranken häufiger an einem Lungenkarzinom als Nichtraucher. Die Ursache dieser Erkrankung ist durch ätiologische Studien als Folge des Tabakkonsums geklärt. Damit ist auch eine Fraktion der Gesamtinzidenz erklärt. Wie hoch dieser Anteil wird durch vergleichende epidemiologische Studien untersucht. Das Ergebnis misst den Inzidenzunterschied zwischen der Studien- und der Vergleichsgruppe. Die Studien hat also den Häufigkeitsunterschied in den gewählten Vergleichsgruppen gemessen. Die Ursache des Unterschieds sind folglich deren Merkmalsunterschiede. Folglich erklärt in diesem Beispiel der Risikofaktor zwar nicht zwingend die Ursache der Erkrankung, wohl aber das Ausmaß des quantitativen Unterschieds zwischen den Vergleichsgruppen.

Die Unterscheidung von Risikomerkmale und Risikofaktor wird immer auf Argumentationen und Hypothesen zurückgreifen müssen, die außerhalb der Interpretationsfähigkeit statistischer Prozeduren liegen.

Unter der hypothetischen Annahme, alle denkbaren Gesundheitsrisiken seien gleichverteilt, bzw. wären vollständig zufallsverteilt, existierten immer noch Risiken, Krankheiten und Krankheitsursachen. Es gäbe allerdings keine Möglichkeit, Risikofaktoren zu ermitteln oder ihrer Auswirkung auf eine Bevölkerung zu quantifizieren.

Die praktische Nutzung von Risikofaktoren zielt auf das prädiktive Argument der Möglichkeit krank zu werden. Da Risikofaktoren aber zumeist auch mit anderen als nur krankmachenden Folgen assoziiert sind, bedarf es regelhaft der Risikobewertung und deren öffentlicher Vermittlung (Risikokommunikation). Da solche Prozesse oft mit erheblichen konkurrierenden individuellen, öffentlichen und wirtschaftlichen Interessen besetzt sein können, kann die Epidemiologie an der Risikobewertung und -kommunikation zwar teilnehmen, sie aber nicht leisten.

Personen mit den gleichen Risikomerkmale/-faktoren werden auch als Risikogruppe bezeichnet.

Da viele Risiken unvermeidlich sind, gehört jeder Mensch mindestens einer Risikogruppe an. Da zudem die Risiken, denen Menschen ausgesetzt sind sich in ihrer Struktur und ihrer Intensität im Laufe des Lebens ändern, lassen sich Menschen immer wieder auch wechselnden Risikogruppen zuweisen. Krankenversicherungssysteme und ebenso Präventionskonzepte, die maßgeblich auf das Konzept der Rasterung von Risikogruppen orientieren, bedürfen deshalb oft aufwändiger Erfassungs- und Klassifikationssysteme, die ggf. ethische und datenschutzrechtliche Akzeptanzprobleme aufwerfen können.

Beispiel:

Drogenkonsum gefährdet die Gesundheit. Diese Feststellung ist für die Epidemiologie relevant, weil sie bspw. zu ermitteln hat, ob die Zahl der Drogenkonsumenten im Ablauf der Zeit und dann in welchen Gruppen der Bevölkerung, zu- oder abnimmt. Aus sozialmedizinischer Perspektive wäre dann zusätzlich zu untersuchen, welche sozialen, bzw. gesellschaftlichen Ursachen und Mechanismen jeweils die Zu- oder die Abnahme der Prävalenz der Drogenkonsumenten erklären. Nachvollziehbar werden hierzu auch andere Wissenschaften, wie z.B. die Psychologie oder auch die Genetik Beiträge zu leisten haben.

2.3 Der epidemiologische Fall

Ein „epidemiologischer Fall“ entsteht durch die Zuordnung (Diagnostik) eines Zustands von Kranksein zu einer durch die ICD bezeichneten Krankheit.

Der Begriff „Fall“ bezeichnet hier einen Zustand von Kranksein an einer definierten Krankheit. Ärzte begegnen in der medizinischen Praxis jedoch keinen Fällen, sondern Individuen. Die abstrakten Krankheitsbegrifflichkeiten der ICD sind also von der individuellen Begegnung von Patienten und Ärzten zu unterscheiden.

Wenn es sich um Krankheiten handelt, die nicht wiederholbar sind, ist die Menge der epidemiologischen Fälle und die Menge der betroffenen Personen gleich groß. Das trifft auch zu, wenn die Fälle nach der Rangordnung ihres Entstehens, etwa in Längsschnittstudien, gezählt werden und die Daten zu den ggf. vielfältigen Kontakte mit den Infrastrukturen zur Krankenversorgung zur Person oder alternativ zu Fällen zusammengeführt werden (können).

Beispiel:

Bei allen wiederholbaren Ereignissen, z.B. die Implantation einer Knieendoprothese kann die Zahl der Behandlungsfälle, die der behandelten verschiedenen Personen übersteigen. Sie können zudem zur Leistungsanspruchnahme in vielen diagnostische, therapeutischen und rehabilitativen Einrichtungen führen.

Für epidemiologische Studien muss immer auch die Qualität der diagnostischen Zuordnung beurteilt werden. Das trifft zwar auch für die Klinik zu, aber mit unterschiedlichen Konsequenzen. Für epidemiologische Studien reicht die Abschätzung der Zuordnungsfehler. Im klinischen Alltag ist alles zu tun, um jeglichen Fehler zu vermeiden.

Beispiel:

Es gibt Schätzungen nach denen der Anteil fehlerhafter Laboruntersuchungen und fehlerhafter Interpretationen solcher Untersuchungen (lab test errors) mit 40 und auch deutlich mehr Prozentpunkten kalkuliert werden muss. Das gilt analog für die bildgebende Diagnostik. Dieser Befund ist aus klinischer und Patientenperspektive dramatisch (Schiff G, Hasan O, Kim S, Abrams R, Cosby K, Lambert B, Elstein AS, Hasler S, Kabongo M, Krosnjar N, Odwazny R, Wisniewski M, McNutt R. Diagnostic error in medicine. Arch Intern Med 2009; 169 (20): 1881-1887). Für epidemiologische Studien kann die Kenntnis dieser Fehleranteile ausreichen sein, um ein Scheitern der Studie zu verhindern.

Bei der falschen Zuordnung von „Fällen“ handelt es sich im Grundsatz immer um Klassifikationsfehler, bzw. um die falsche Zuordnung eines Zustandes zu einem Klassifikationsbegriff. Bei der Beurteilung der Qualität der Zuordnung epidemiologischer Fälle sind vier Sachverhalte zu unterscheiden:

1. die Spezifität
2. die Sensitivität
3. der positive Vorhersagewert (positive predictive value)
4. der negative Vorhersagewert (negative predictive value)

Die Kenntnis dieser Klassifikationsfehler gehört zum ärztlichen Basiswissen (siehe hierzu Galen, R.S., Gambino, S. R. (1976): *Beyond Normality: The Predictive Value and Efficiency of Medical Diagnoses*. Churchill Livingstone New York, 1976 ISBN 10 0471290475).

Je geringer die Wahrscheinlichkeit ist, dass unter einer Menge Untersuchter auch Kranke gefunden werden können, desto geringer ist der Aussagewert des diagnostischen Verfahrens.

Das hat für epidemiologische Studien, die der interventionellen Epidemiologie zugeordnet werden (Interventionsstudien) und Screeningmethoden einsetzen, eine nennenswerte Bedeutung. Gemäß dem Allgemeinen Krankheitsmodell (siehe Kapitel 4.1) ist deren Ziel u.a. die Verkürzung der diagnostischen Intervalle und zeitliche Vorverlagerung der Klassifikation in Zustände des Krankseins sowie auch um die Diagnostik von Zuständen des Gefährdetseins und um medizinischen Interventionen in Vor- und Frühformen von Kranksein. Die quantitativen Folgen für die resultierenden neuen Bedarfe sind erheblich und übersteigen die Bedarfe ohne solche prädiktiven Fahndungen deutlich. Diese Folgen für die Prävalenzen ändern jedoch immer auch die prädiktiven Werte der diagnostischen Tests. Die Anzahl der falsch positiven kann dann

leicht, die der richtig positiven Befunde übersteigen. Es ist deshalb zwischen dem Einsatz diagnostischer Methoden unter den Bedingungen medizinischer Versorgung, also im Krankenhaus oder in der Arztpraxis und dem Einsatz unter den Bedingungen einer systematischen Früherkennung (Screening) etwa im Zusammenhang mit epidemiologischen Studien und antiepidemischen Konzepten zu unterscheiden.

Die Interpretation diagnostischer Ergebnisse in der Früherkennung und unter den Bedingungen des klinischen Alltags sind grundsätzlich zu unterscheiden.

Diagnostikangebote in der ärztlichen Praxis zur Früherkennung sind nicht mit einem Screening zu verwechseln. Solche Angebote sind oftmals schlicht ein Stochern im Nebel und entziehen sich der empirisch-analytischen Rationalität.

Bei der Bewertung diagnostischer Klassifikatoren sind folgende Fragen zu beantworten:

- Wie ist das quantitative Verhältnis zwischen Symptomträgern (bzw. Personen mit Befunden, die außerhalb der Norm liegen) und Kranken?
- Mit welcher Wahrscheinlichkeit ist jemand krank (bzw. gesund), der ein positives (bzw. negatives) Testergebnis hat?
- Wie viele der tatsächlich Kranken (Gesunden) werden mit dem Test auch tatsächlich als krank (gesund) richtig erkannt?

Testresultat	krank		Total
	ja	nein	
positiv	a	b	a + b
negativ	c	d	c + d
total	a + c	b + d	a + b + c + d

Tab. 1: Vierfeldertafel für den Zusammenhang von Krankheit und Testresultat

Die hierbei zu beachtenden Zusammenhänge lassen sich (im Falle dichotomer oder binärer Entscheidungen) an einer **Vierfeldertafel** darstellen.

Die aus dieser Tafel ableitbaren Ziffern beantworten jeweils spezifische Fragen.

Bezeichnung	Die zuzuordnende Frage
Sensitivität	Wie viele von den Kranken werden mit dem Test als krank erkannt?

Spezifität	Wie viele von den Gesunden werden mit dem Test als gesund erkannt?
positiver prädiktiver Wert	Mit welcher Wahrscheinlichkeit ist jemand mit einem positiven Testergebnis krank?
negativer prädiktiver Wert	Mit welcher Wahrscheinlichkeit ist jemand mit einem negativen Testergebnis gesund?

Tab. 2: Gütekriterien eines diagnostischen Tests

2.4 Früherkennung

Früherkennung unterscheidet die Diagnostik früher Zustände von Kranksein und die Diagnostik von Zuständen, die mit einer zu messenden Wahrscheinlichkeit von Krankwerden gefolgt sind (Risikodiagnostik).

Neben der Bekämpfung von Epidemien sind die Früherkennung von Krankheiten, ggf. bereits von Gefährdungen und individuellen Dispositionen in vielen Staaten gesundheitspolitische Schlüsselziele des Public Health und legale Pflichten (siehe hierzu Heft 8). Diese Ziele verlangen dann auch eine kapazitiv und methodisch leistungsfähige Epidemiologie.

Bei der Einführung von Früherkennungsprogrammen ergibt sich das Dilemma, dass eine Entscheidung zwischen dem Wunsch nach der möglichst vollständigen Erkennung aller Kranken bzw. Gefährdeten und einer möglichst geringen Anzahl fälschlich als krank bzw. krankheitsgefährdet Klassifizierten herbeigeführt werden muss. Ein guter Früherkennungstest soll bei möglichst wenigen Personen mit der Erkrankung und bei möglichst vielen Personen ohne Erkrankung negativ sein. Dabei geht eine Steigerung der Sensitivität des Testes meist mit der Verringerung der Spezifität und des positiven prädiktiven Wertes des Testes einher.

Da eine höhere Sensitivität bei gleicher Untersuchungsmethodik oft nur durch Verschiebung des Normwertbereiches realisierbar ist, folgt zwangsläufig eine wachsende Anzahl falsch positiver Testresultate (Gesunde mit einem positiven Testresultat). D.h., die Steigerung der Sensitivität eines Tests durch Verschiebung des Normwertbereiches führt zu einer Verringerung der Spezifität und des positiven prädiktiven Wertes dieses Testes. Früherkennungen sind, von anderen Kriterien abgesehen, bei Gefährdungen und frühen Formen einer Krankheit mit einer niedrigen Prävalenz folglich wenig sinnvoll.

Vortestwahrscheinlichkeit	LR+ = 19 Sensitivität = 95 %, Spezifität = 95 %		LR+ = 2.3 Sensitivität = 70 %, Spezifität = 70 %	
Prävalenz	PV+	PV-	PV+	PV-

90 %	99 %	67 %	96 %	21 %
50 %	95 %	95 %	70 %	70 %
10 %	67 %	99 %	21 %	95 %
1 %	16 %	100 %	2 %	97 %

Tab. 3: Zusammenhang von Prädiktion (PV), Likelihood-Ratio (LR) und Prävalenz (nach B. Höldke).

„Da auch Personen ohne Erkrankung ein positives Testergebnis aufweisen können, ist es wichtig abzuschätzen, um wieviel häufiger ein positives Testresultat bei Personen mit Erkrankung im Vergleich zu Personen ohne Erkrankung vorkommt. Dieses Verhältnis wird Likelihood Ratio genannt. Da solche Schätzungen nur in epidemiologischen Studienkontexten möglich sind, ist die Effektivität von Früherkennungen im Zusammenhang mit medizinischen Behandlungsanlässen kritisch zu befragen.“

2.5 Epidemiologie und Ätiologieforschung

Forschungen zur Ätiologie fragen nach den Ursachen, die am Beginn pathogener Prozesse standen, diese also gleichsam in Gang gesetzt haben. Auch alle nachfolgenden einzelnen und komplexen pathogenetischen Vorgänge haben selbstverständlich Ursachen, sind also dem Kausalitätstheorem unterworfen.

Die Frage, welchen Beitrag die Epidemiologie zur Aufklärung der Ätiologie von Krankheiten leisten kann, bedarf einer differenzierten Antwort. Es ist zu unterscheiden, ob nach der Ursache des Krankwerdens einer konkreten Person oder nach allen Ursachen aller Fälle von Krankheit gleicher Diagnose gefragt wird.

Beispiel:

Das Kind Peter erleidet beim Sturz vom Fahrrad eine Fraktur (nach ICD als V99 zu klassifizieren). Würden die Eltern von Peter wegen dieses Sturzes haftungsrechtliche Konsequenzen anstreben, wäre durch einen Gutachter die „Ätiologie“ des Unfalls zu klären. Das Gutachten diene dann der Beweissicherung in einem Einzelfall, bzw. in Bezug auf eine konkrete Ursache-Wirkungsbeziehung. Aus der Perspektive einer epidemiologischen Fragestellung ist das Szenario völlig anders: Nachdem eruiert ist wie viele Fahrradstürze mit Frakturfolge es in einer Bevölkerung in einem Beobachtungszeitraum gibt, wäre der nächste Schritt, die Struktur der Unfallopfer (z.B. Alter und Geschlecht, Nutzungsart des Fahrrads, Unfallszenario, Randbedingungen wie Tageszeit und Wetter etc.) zu ermitteln. Das Ergebnis wäre, dass es viele verschiedene Ursachen für Fahrradstürze mit Frakturfolge gibt. In Bezug auf Peter und die Haftungsfrage seiner Eltern ist die Argumentation, die unter V99 zusammengefassten Fälle seien multifaktoriell verursacht, unsinnig. Die epidemiologische Analyse würde allerdings Verteilungskennntnisse verfügbar machen, die beschreiben, welche Kinder im Alter von Peter solche Unfälle eher

erleiden als andere und welche Wertigkeit dabei (fraktionell) einzelne Ursachen haben. Möglich wären Unterschiede bezüglich der technischen Sicherheit oder bezüglich von Ursachen, die in der Person der Unfallopfer oder anderer Unfallverursacher liegen. Die Bewertung solcher Ursachenstrukturen könnte dann Prioritäten für die Prävention setzen, sind aber in der haftungsrechtlichen Fragestellung in der Regel nicht relevant.

Das Potenzial der Epidemiologie gründet auf ihren Möglichkeiten, die Struktur der Ursächlichkeiten zu ermitteln. Die Epidemiologie kann dann ggf. auch den Strukturwandel der Ursachen beschreiben und auch diesen Wandel ursächlich klären. Sie kann keine ursächlichen Entscheidungen im Einzelfall treffen.

Die Epidemiologie macht keine Aussagen zum Krankwerden einzelner Individuen.

Beispiel:

In längerdauernden schweren sozialen und wirtschaftlichen Krisenzeiten nimmt die Anzahl der Personen mit Neuerkrankungen an Tuberkulose typischerweise zu, die der Personen mit einem Diabetes Typ 2 hingegen ab. Beide dieser Veränderungen gehen ursächlich auf die soziale Krise zurück. Im Falle der Tuberkulose ist dies ggf. einer höheren Anfälligkeit zuzurechnen, im Falle des Diabetes Typ 2 einer Verringerung der Krankheitsdauer und/oder einer Verringerung der Diagnosestellungen infolge des Zusammenbruchs der medizinischen Versorgung.

Da die Epidemiologie stets nur Aussagen zu „Mengenveränderungen“ und deren Ursächlichkeit machen kann, sind ihre Feststellungen auch nur auf diese, nicht auf einzelne Personen zu beziehen. Warum also eine konkrete Person an Tuberkulose erkrankt oder eine andere nur (noch) eine kurze Behandlungsdauer erfahren hat, kann hierdurch ggf. wahrscheinlich gemacht, nicht aber für den Einzelfall abschließend geklärt werden.

In der Regel kann davon ausgegangen werden, dass Ursachen epidemiologisch relevanter Verteilungsänderungen die Gesamtheit der Personen einer Bevölkerung nie in gleichem Ausmaß betreffen. Folglich werden solche Veränderungen zu Häufigkeitsunterschieden zwischen einzelnen Teilbevölkerungen führen.

Ursachen erklären Wirkungen. Folglich kann es ohne eindeutige Bestimmung der zu klärenden Wirkung auch keine ursächlichen Zuordnungen geben. Eindeutige Ursache-

Wirkungs-Beziehungen lassen sich immer nur dann rekonstruieren, wenn die Wirkung bereits vorliegt, also z.B. ein Individuum bereits krank ist (ex post). Ursachen sind retrospektiv und in Bezug auf eine konkrete Wirkung im Einzelfall definiert. Dieser Einzelfall kann verallgemeinert werden, wenn jeweils gleiche oder zumindest ähnliche Randbedingungen anzunehmen sind. Prospektive Ursachenannahmen sind hingegen Risiken oder Gefährdungen (ex ante).

Beispiel:

Die Anamnese des Patienten A legt die Schlussfolgerung nahe, dass dessen Lebererkrankung die Folge langdauernden Alkoholmissbrauchs ist. Die gleiche Lebererkrankung des Patienten B ist hingegen ursächlich als Folge langdauernder toxischer Einwirkungen am Arbeitsplatz zu beurteilen. Beide Erklärungen verweisen auf eine unterschiedliche Ätiologie eines ähnlichen Krankheitsbildes, nicht auf eine multifaktorielle Genese einer Lebererkrankung.

Beispiel:

Bei einem Unfallopfer wird ein massiver Blutdruckabfall festgestellt. Das ärztliche Interesse an der Kausalität ist erfüllt, wenn als Ursache eine Milzruptur diagnostiziert wird. Im Sinne des Handlungsziels, nämlich die Sicherung vitaler Funktionen, ist die Feststellung dieser Ursache hinreichend. Wenn diese Ruptur jedoch die Folge einer Gewalthandlung gewesen ist, ist aus der Sicht staatsanwaltlicher Ermittlungen die Blutung die Folge eben dieser Gewalthandlung. Deren strafrechtliche Beurteilung bedarf wiederum eigener ursächlicher Erklärungen.

Beispiel:

Die Zunahme der mittleren Lebensdauer erklärt ursächlich den Wandel der Struktur der Todesursachen. Das bedeutet selbstverständlich nicht, dass die Krebserkrankung einer konkreten Person und ggf. der hieraus folgende Tod ursächlich durch die Zunahme der mittleren Lebensdauer erklärt wäre.

Die Identifikation von Ursachen ist zweckgebunden. Für unterschiedliche Ziele können sehr unterschiedliche Ursachenzusammenhänge relevant sein.

Ursachen sind immer nur als Beziehungen zu konkret definierten Wirkungen zu bestimmen. Jede Änderung der Definition einer untersuchten Wirkung wird auch die ihr zurechenbare Ursache verändern. Je komplexer Wirkungen definiert werden, desto

weniger eindeutig ist die zurechenbare Ursache. Ein spezieller Aspekt berührt deshalb die Frage, ob die Epidemiologie für gutachtliche Fragestellungen etwa bei krankheitsbegründeten Haftungsansprüchen im Einzelfall als Beweismittel herangezogen werden kann. Diese Frage ist im Grundsatz zu verneinen. Die Fragestellungen ist jedoch völlig anders, wenn zu beurteilen ist, wie viele Personen Opfer einer Werbekampagne für gesundheitsschädigende Produkte geworden sind (attributable risk fraction). Die Einzelfälle der Opfer lassen sich durch epidemiologische Analysen nicht identifizieren. Wie mit solchen Problemen rechtlich umzugehen ist, kann nur im Kontext der jeweiligen nationalen Rechtslagen entschieden werden.

Beispiel:

Herr A. klagt, nach dem er Kenntnis von einer Karzinomerkrankung erhalten hat, auf Haftung, weil der Tabakkonzern X ihm beim Kauf von Tabak keine Warnung vor den Gefahren des Rauchens mitgegeben hat. Der medizinische Gutachter steht hier ggf. vor dem Dilemma, erklären zu müssen, ob die Erkrankung auf den Tabakkonsum zurückzuführen ist oder nicht. Dieser Nachweis wird ihm im Einzelfall kaum mit letzter Sicherheit gelingen. Wenn hingegen eine solche Haftungsklage in einem Rechtssystem verhandelt würde, das Sammelklagen erlaubt, stünde die Frage in einem völlig anderen Kontext. Der Gutachter könnte darauf verweisen, dass unter Beachtung der Konfidenzintervalle die a priori Wahrscheinlichkeit für die Entstehung eines Lungenkarzinoms infolge des Rauchens bei einer durchschnittlichen Konsumdauer x und einer verbrauchten Tabakmenge y z.B. 10% beträgt. Unter 1.000 klagenden Personen mit einem Lungenkarzinom und einer vergleichbaren Anamnese dürfte also für 100 Kranke angenommen werden, die Krebserkrankung sei Folge des Rauchens. Es wäre im Falle zulässiger Sammelklagen hier bedeutungslos, dass nicht entschieden werden kann, welche konkrete Person wegen des Rauchens oder aus anderen Gründen erkrankte. Es könnte entschieden werden, dass für 100 Personen eine Entschädigung zu zahlen ist, die ggf. unter allen 1.000 Klagenden aufzuteilen wäre. Diese Fragestellung gewänne noch eine andere Dimension, wenn nach einer nationalen Rechtslage alle Raucher im Falle einer Lungenkrebserkrankung von einem Anspruch auf Leistungen ihrer Krankenversicherung ausgeschlossen würden, weil unterstellt wird, sie hätten ihre Erkrankung selbst verschuldet. Bezogen auf das gegebene Beispiel würde u.U. 90% der Kranken ungerechtfertigt eine medizinische Hilfeleistung verwehrt.

In dem Beispiel ist die Richtigkeit der Aussage, Rauchen kann Lungenkrebs erzeugen, nicht berührt. Die Aussage, ein konkretes Individuum, das an einem Lungenkarzinom leidet, habe dieses wegen seines Tabakkonsums selbst verschuldet, ist hingegen unzulässig.

An der Schnittstelle zwischen qualitativen und quantitativen Urteilen existiert folglich ein Spannungsfeld. Dieses Spannungsfeld ist über das Beispiel hinaus von grundsätzlicher Bedeutung. Es betrifft die Frage, ob und unter welchen begrenzenden Bedingungen stochastisch begründete Aussagen für den Einzelfall akzeptiert werden

können. Bei solchen Entscheidungen kann es sich um präventive Empfehlungen ebenso handeln wie um therapeutische oder gutachtliche. Allerdings wird ihre Konsequenz jeweils erheblich verschieden sein. Die Konsequenz ist bei der Empfehlung aus Gründen der Prävention nicht zu rauchen, unproblematisch, im Falle einer Therapieverweigerung für Raucher kann sie hingegen schwer zu rechtfertigen, ggf. rechtswidrig sein.

Während der Ursachenbegriff retrospektiv (entsprechende Datenlage und entsprechenden Erkenntnisstand vorausgesetzt) zu eindeutigen Erkenntnissen führt, lassen sich in der Umkehrung in der Regel prospektiv keine eindeutigen Voraussagen (Prädiktionen) für einzelne Individuen herleiten. Es können jeweils nur Aussagen über die Wahrscheinlichkeit eines Folgeereignisses gemacht werden. Diese kann aber im Einzelfall immer nur 0 oder 1 sein.

Diese Schnittstellen können auch so charakterisiert werden:

- Epidemiologie ist leistungsfähig, wenn abgeschätzt werden soll, wie hoch der Anteil jener Personen ist, deren Krankwerden jeweils spezifischen Ursachen zuzurechnen ist. Hierbei ist die Korrektheit der individuellen Zuweisung unerheblich (Prädiktion von Mengen).
- Epidemiologische Aussagen sind eine problematische Entscheidungsgrundlage, wenn individuelle Beurteilungen und Empfehlungen (haftungsrechtliche Gutachten, prospektive Entscheidungen über therapeutische Maßnahmen im Krankheitsfall etc.) angestrebt werden (Prädiktion von Einzelereignissen).

Wahrscheinlichkeiten können also für eine konkrete Person nur mit erheblicher Unsicherheit angegeben werden. Abgeschätzt werden kann allerdings, wie viele Personen unter gleichen Bedingungen erkranken werden. In der Praxis wird dann die quantitative epidemiologische Analyse durch eine immer mehr oder weniger auch subjektive qualitative Deutung zu ergänzen sein.

Risikobewertungen münden alltagssprachlich in die Begriffe Chance und Risiko. Diese sind Bewertungen der Wahrscheinlichkeit des Auftretens des Zielereignisses. Da solche Bewertungen kontextgebunden sind und ggf. von komplexen konkurrierenden Interessen abhängig sein können, können Epidemiologen an solchen Bewertungen zwar mitwirken, sie aber nicht allein übernehmen. Es gibt im Einzelfall keinen epidemiologischen Maßstab für die Unterscheidung von Risiko und Chance.

Risikobewertungen sind kontextgebunden und unterliegen konkurrierenden Interessen. Risikobewertungen sollten deshalb möglichst weitgehend Gegenstand öffentlicher Konsensbildung sein. Hieran nimmt die Epidemiologie beratend teil.

2.6 Erneuerungsmengen

Epidemiologische Potenziale sind Erneuerungsmengen. Erneuerungsmengen erneuern sich durch Zugänge und Abgänge in der Verweildauer in ihren epidemiologischen Potentialen.

Ohne die Erneuerung epidemiologischer Potenziale ist deren „Bewegungen“ nicht möglich. Das Verständnis solcher Erneuerungen ist folglich für die Epidemiologie essenziell.

Die Erneuerungsdauern entscheiden maßgeblich über die Dynamik der jeweiligen Prävalenzen, bzw. über die Dynamiken von Epidemien und Regressionen.

Bei den Personen, die unterscheidbaren Potenzialen zugeordnet werden können, kann davon ausgegangen werden, dass sie sich in ihren Lebensumständen und in ihren Vulnerabilitäten ähneln, wenngleich sie nie identisch sind. Diese Potenziale entsprechen den jeweiligen Gruppen von unterschiedlich gefährdeten Personen. Sie unterscheiden sich von andern durch ihre quantitativen und strukturellen Eigenschaften. Sie sind z.B. durch ihre skalierbare spezifische Anfälligkeit (Disposition), durch gemeinsame Expositionen und Expositionsdauern, durch ihre spezifischen kollektiven Präventionsinteressen und ihr präventives Verhalten, durch Alter, Geschlecht, Komorbiditäten, ihre allgemeine Lebenssituation etc. unterschieden.

Die Erneuerungsprozesse dieser Potenziale sind ein analytisches Kernproblem der Epidemiologie. Sie sind durch zwei Aspekte zu charakterisieren. Ein erster besteht darin, dass diese Erneuerungen sowohl durch die Veränderungen der Mengen neuer Fälle (Inzidenz), als auch durch die Veränderungen der Verweildauer in den jeweiligen Potenzialen bewirkt sein können. Hieraus folgen dann die für deren Charakteristik entscheidenden Mengenänderungen der Potenziale und verändert so die Verläufe von Epidemien oder Regressionen. Für den Ablauf ihrer Zyklen ist der mit den Mengenänderungen verbundene Strukturwandel der Prävalenzen ein entscheidender Vorgang. Neben der Bewegung der Mengen, ist also der Strukturwandel für die entsprechenden Bewegungen, verursacht durch die jeweiligen Bewegungsgrößen, der entscheidende Vorgang.

Es ist deshalb gerechtfertigt, den Strukturwandel der Potenziale, also die epidemiologischen Potenziale der Gefährdeten, der Kranken und der aus der Menge der Kranken Ausscheidenden als die ihre jeweiligen Bewegungen maßgeblich formenden Größen zu bezeichnen. Diese Strukturveränderungen haben vor allem folgende Ursachen:

- den „Verbrauch“ der Potenziale in der Rangfolge ihrer jeweils spezifischen Vulnerabilitäten
- das unterschiedliche Erlernen von präventiven Strategien
- die sozial varianten Bereitschaften und Möglichkeiten zur Prävention und

- die sich verändernden therapeutischen Möglichkeiten und die ggf. soziale Varianz deren Zugänglichkeiten

Diese Änderungen betreffen zumeist die Strukturmerkmale Alter, Geschlecht, Begleitkrankheiten, soziale Schichtzugehörigkeit und individuelle Ressourcen. Das bedeutet aber auch, dass sich während einer Epidemie/Regression die präventiven und therapeutischen Konzepte sowie die zur Krankheitsfolgenbewältigung sich immer wieder an veränderte Bedingungen anpassen müssen.

2.7 Die Reproduktionsrate

Die Reproduktionsrate bilanziert die Bewegungsgrößen der Prävalenzen.

Der erste uns bekannte Versuch, die Bewegung einer Prävalenz mit mathematischen Mitteln zu beschreiben bezog sich auf eine übertragbare Krankheit, auf die Malaria, bzw. die Mengenbewegung ihrer Verursacher Anopheles. Hierzu beschrieb der Entdecker des Malariaparasiten im Jahr 1897, der Physiologen und Nobelpreisträger des Jahres 1902, Ronald Ross (1857-1932) die Basisreproduktionsziffer (BRZ). Die BRN sollte anzeigen, ob sich die neu auftretenden Fälle (Inzidenz) und die existierenden Fälle (Prävalenz) verlassenden Fälle in einem Gleichgewicht befinden oder eben nicht.

Die BRZ war offenbar der erste Versuch, eine Epidemie als ein Ungleichgewicht der Prävalenz zu deuten. Damit definierte Ross zwangsläufig auch die Regressionen der Prävalenz. Die von Ross entwickelte Reproduktionszahl wurde später von Gregor Macdonald weiterentwickelt (Macdonald, G. (1957): *The Epidemiology and Control of Malaria*. Oxford University Press, London; siehe auch Smith, DL, et al "Ross, Macdonald, and a theory of the dynamics and control of Mosquito-transmitted pathogens". *PLOS Pathogens*. 8 (4): 1002588. Doi: 10.1371/journal.ppat.1002588. ISSN 1553-7366). Der Begriff „Reproduktionsziffer“ steht heute im Kontext der Epidemiologie für mehrere mathematische Modelle und bestimmt eine Epidemie/Regression grundsätzlich als die Bewegung von Prävalenzen.

Die BRN wurde zu einer Zeit entwickelt, als noch niemand die Frage diskutierte, ob Epidemien auch solche Arten von Falldynamiken umfassen, die durch andere Gründe als die Infektionsdynamiken verursacht werden (siehe hierzu auch Ross, R. (1905) *The logical basis of the sanitary policy of mosquito reduction*. *Science* 22: 689-699. 22, oder Ross, R.R. (1911): *Some quantitative studies of epidemiology*. *Nature* 87: 466-467. 25). Macdonald beschrieb dann die Tatsache, dass die Anzahl der Fälle, die auf einen ersten Fall folgen, auch von dem Intervall abhängt, in dem eine Person oder ein Wirt andere infizieren kann. Die Variation und Veränderung dieses Intervalls hat also Auswirkungen auf diese Mengen und ihre Ungleichgewichte.

Wenn die Fälle, die die Prävalenz innerhalb des Zeitintervalls, in dem sie infektiös oder krank waren, verlassen, die gleichen Beträge haben wie die inzidenten Fälle, ist die Reproduktionszahl der Prävalenz immer 1. Die Prävalenz befindet sich unter dieser

Bedingung im Zustand der Stationarität. Das gilt selbstverständlich für alle Prävalenzen Gefährdeter, in Falle von Infektionskrankheiten für die Infizierten, dann für die Kranken und alle weiteren relevanten Bestandsmengen.

Die Konsequenz ist fundamental: Jede Nicht-Stationarität muss andere Ursachen haben, als die Eigenschaften der Exposition erklären können. Im Sonderfall einer Malaria-Epidemie musste diese also durch die Ursache erklärt werden, die die Reproduktionsmuster von Anopheles verändert. Diese Ursache ist nicht Plasmodium vivax selbst. Schon John Snow (1813-1858) konnte zeigen, dass der Vektor eine Epidemie, in diesem Falle einer Choleraepidemie, nicht aber der ätiologische Faktor, der die Krankheit verursacht, die Epidemie in Gang setzte und dann auch unterhielt.

Wir wissen heute, dass dies für jede Epidemie und jede Regression jeglichen epidemiologischen Potenzials gilt, also nicht nur für Epidemien, die durch biologische Expositionen bedingt sind. Und genau das macht die Epidemiologie zur Wissenschaft von den Dynamiken der Prävalenzen und nicht zur Wissenschaft von der Verursachung (Ätiologie) bestimmter Krankheiten.

Die "Basisreproduktionszahl" muss von der "effektiven Reproduktionszahl" unterschieden werden. Diese anerkennt, dass sich im Ablauf einer Epidemie die Struktur des gefährdeten epidemiologischen Potenzials verändert. Der Strukturwandel folgt nicht aus einer veränderten Ätiologie, sondern aus den Veränderungen der Potenzialeigenschaften, also etwa denen der sozialen Bedingungen, der Einstellungen, der Bildung oder des Verhaltens der Menschen, im Falle übertragbarer Krankheiten, des immunologischen Status des zugehörigen Potenzials.

Die Reproduktionszahlen sind weder in der Lage, eine neue Epidemie vorherzusagen, noch können sie zur Ursachenanalyse herangezogen werden, warum sich die Mengen der prävalenten Fälle ändern. Die das mathematische Modell begründenden Annahmen sind für Public Health Maßnahmen jedoch von praktischem Nutzen, wenn Epidemiologen in der Lage sind, diese Zahl für die Vielfalt der relevanten epidemiologischen Potenziale differenziert zu ermitteln, also deren soziale und regionale Muster und deren Veränderungen im Zeitverlauf abzuschätzen. Das Problem im Umgang mit der BRZ ist, dass das durch sie definierte Gleichgewicht von der Größe der Prävalenz unabhängig ist.

2.8 Diagnostische und therapeutische Intervalle

Das diagnostische Intervall ist der Zeitraum zwischen dem Auftreten einer Krankheit und der Feststellung der Diagnose. Das therapeutische Intervall ist der Zeitraum zwischen der Diagnose und dem Ende der therapeutischen Maßnahmen.

Diese Intervalle sind für alle Krankheiten, wenngleich in unterschiedlichem Ausmaß, bedeutsam. Der Grund liegt in ihrer Ungleichverteilung innerhalb eines individuellen Lebens sowie auch in den Gesamtbevölkerungen, in gegeneinander zu unterscheidenden epidemiologischen Potenzialen und schließlich auch in der Veränderbarkeit

der Intervalle im Ablauf der Zeit. Hieraus folgen dann auch quantitative Veränderungen, die sowohl Epidemien wie auch Regression vortäuschen können.

Bei den übertragbaren Krankheiten (Bakterien, Viren, Parasiten) gibt es ein zusätzliches Intervall. Dies ist der Zeitraum zwischen einer Infektion und der Wahrscheinlichkeiten, dass einer Infektion auch eine Erkrankung folgt. Hier sind die infizierten jedoch nicht kranken Personen ein eigenes epidemiologisch bedeutsames Potenzial, weil nicht erkrankte, aber infizierter Personen hochrelevante Infektionsquellen sein können.

Bei akut auftretenden und schweren Krankheiten, wie z. B. übertragbaren Krankheiten, sind die Diagnoseintervalle in der Regel kürzer als bei nicht übertragbaren Krankheiten, wenn auf eine Infektion sofort Symptome folgen. Bei solchen Krankheiten sind auch die therapeutischen Intervalle oft hinreichend konstant. Aus diesem Grunde können dann die Inzidenzen als Anhaltswerte für die Dynamik der Prävalenzen herangezogen werden.

Die Dauerverteilungen dieser Intervalle sind in der Regel asymmetrisch oder sogar multimodal. Sie spiegeln keine Krankheitseigenschaften, sondern vor allem die Qualität der medizinischen Versorgungssysteme und deren Nutzung wider.

Die Verteilung dieser Intervalle innerhalb einer Population ändert sich in der Regel während einer Epidemie/Regression mit bemerkenswerten quantitativen Auswirkungen. Solche zeitlichen Auswirkungen sind für das Verständnis der Dynamik der Fallzahlen und deren strukturelle Merkmale wie Alter, Geschlecht, Schweregrad oder Heilungsraten und sozialer Status bedeutsam.

Analoge Modifikatoren wirken auf die therapeutischen Intervalle, die entsprechend kürzer oder länger werden können und so die Verweildauern im Bestand der Kranken und folglich auch für die Ausgänge von Kranksein erhebliche quantitative Folgen haben.

2.9 Die epidemiologische Durchdringung

Die epidemiologische Durchdringung (Penetration) beschreibt die Mechanismen, durch welche sich die Expositionen auf die verschiedenen epidemiologischen Potenziale verteilen.

Von Expositionen abhängige Erkrankungen benötigen spezifischer Mechanismen, um sich in ihren relevanten epidemiologischen Potenzialen verteilen zu können. Diese Mechanismen sind durch die bevorzugten "Vektoren" des Eindringens zu beschreiben, die in physikalische, biologische oder soziale Vektoren unterschieden werden können. Das Eindringen der Exposition in eine Bevölkerung ist dann der treibende Mechanismus einer Epidemie und durch ihre zeitlichen Parameter zu beschreiben. Analog ist deren „Verdrängung“ dann ein Mechanismus, der eine Regression erklärt.

Die Wirksamkeit des Eindringens einer Exposition in eine Bevölkerung hängt von den Eigenschaften der betroffenen Bevölkerungen ab. Solche Merkmale sind

1. die Anzahl der effektiven Kontakte zwischen Exposition und Potenzial
2. die Anfälligkeit gegenüber einer für Exposition
3. die Intensität (Stärke) einer Exposition

Diese Merkmale bestimmen die spezifischen Eigenschaften einer Epidemie/Regression. Und jedes dieser Merkmale muss durch deren Verteilung auf die verschiedenen epidemiologischen Potenziale beschrieben werden. Obwohl es empirisch schwierig sein mag, die Anfälligkeit einer Bevölkerung zu messen und zu klassifizieren und die Intensität der infektiösen Exposition zu kalkulieren, ist dies grundsätzlich möglich.

Während Anfälligkeit, synonym Vulnerabilität oder Empfänglichkeit, die Struktur der Gefährdeten beschreibt, wie z.B. mangelnde Immunität oder das Leiden an Komorbiditäten, kann diese auch unter dem „Sense of Coherence“, wie A. Antonovsky (1923-1994) ihn beschrieben hat, zusammengefasst werden (siehe Antonovsky, A.: *Unraveling the Mystery of Health*). Jossey Bass, San Francisco 1987).

Beispiel:

Einer der von Epidemiologen oft zitierten historischen Unfälle war der Untergang der Titanic. Analysen haben gezeigt, dass die Überlebenswahrscheinlichkeit von dem gemieteten Deck abhing. Dies war jedoch eindeutig von der sozialen Schicht abhängig.

Die Wirksamkeit von Präventionsmaßnahmen gegen von Gefährdungszunahmen verursachte Epidemien geht zunächst auf Interventionen zur Verringerung der Verbreitung von Expositionen zurück. Dies verdeutlicht, dass Seuchenprävention eine gesellschaftliche Maßnahme ist, die immer auch die gegebenen Bedingungen des täglichen Lebens berücksichtigen muss. Es gehört heute zum Grundwissen, dass die erfolgreichsten Präventionserfolge gegen Epidemien lange vor der Verfügbarkeit wirksamer medizinischer Interventionen wie z.B. der Impfung, Realität wurden. Das hat sich im Verlauf der Geschichte verändert. Es gibt Beispiel, nach denen therapeutisch Konzepte offenbar wirksamer und auch effizienter sind als präventive.

Die Begrenzung der Durchdringung der Bevölkerung erfordert längerfristige soziostrukturelle Veränderungen der Arbeits- und Lebensbedingungen sowie eine Verbesserung der allgemeinen Bildung. Eine rasche Begrenzung der Exposition wird ohne gesetzliche Regelungen nur schwer zu erreichen sein. Dies macht deutlich, dass die erfolgreiche Bekämpfung von Epidemien immer den politischen Willen und die Fähigkeit zum Management von Gesundheitsrisiken voraussetzt.

2.10 Die Intensität von Expositionen

Die Intensitäten von Expositionen sind deren spezifische Eigenschaften.
--

Intensitätsmessungen sind für die Bewertung von Expositionen ebenso wichtig wie z.B. für die Bewertung von pharmazeutischen Wirkstoffen. Epidemiologen verwenden in der Regel Intensitätsmessungen unter experimentellen Bedingungen.

Beispiel:

Der Biss einer Klapperschlange ist für den Menschen tödlich, und ihr Gift gehört wohl zu den größten Gefahren, die wir kennen. Aber selbst in den Teilen der Welt in denen diese Schlangen leben, ist diese Gefahr als Problem der öffentlichen Gesundheit sehr unbedeutend. Ganz anders verhält es sich mit Sars-CoV-2 oder Vibrio cholerae. Nicht nur die Intensität dieser infektiösen Expositionen ist hoch. Hinzu kommt die Wahrscheinlichkeit, unter bestimmten Lebensbedingungen mit diesem Virus in Kontakt zu kommen. Folglich ist die Verhinderung des Eindringens dieses Virus in eine Bevölkerung und innerhalb einer Bevölkerung eine Herausforderung von hoher Priorität für die öffentliche Gesundheit.

Für die Epidemiologie bleibt vor allem die Aufgabe die Relevanz, also die Durchdringung einer Bevölkerung von Expositionen zu ermitteln (siehe Kapitel 5).

2.11 Die Zeit unter Risiko

Die Einwirkungsdauer von Expositionen, auch als „Zeit unter Risiko“ bezeichnet, entspricht der Zeitspanne zwischen dem Beginn einer Exposition eines epidemiologischen Potenzials und den Ereignissen nachfolgender Erkrankungen.

Diese Dauer bezieht sich also auf Personen, die in Bezug auf eine definierte Exposition und ihre potenzielle Folgekrankheit (noch) nicht erkrankt sind. Sie ist für die Erkrankenden nur ex post zu ermitteln und dient zur Kalkulation der Erkrankungswahrscheinlichkeit als Funktion der Zeit. Diese Dauer kann zwischen den bei Menschen vorkommenden Krankheiten und zwischen den Betroffenen erheblich variieren. Für die Epidemiologie ist von besonderer Bedeutung, dass die Verteilung dieser Zeitspanne zwischen disjunkten epidemiologischen Potenzialen variieren kann, Sie ist deshalb für diese Potenziale ggf. definitionsrelevant.

Menschen sind ihr ganzes Leben lang dem Risiko ausgesetzt, krank zu werden. In dieser Perspektive nimmt deshalb diese Zeitspanne mit der erreichten Lebensdauer zu und zugleich ändert sich auch die Struktur der Gefährdeten und die der relevanten Expositionen.

Die Zeit unter Risiko kann das ganze Leben umfassen und auch Generationen übersteigen. Aus praktischen Erwägungen ist es deshalb sinnvoll, solche Intervalle in Bezug auf eine genau definierte Gefahr und ihre spezifischen Auswirkungen auf die Gesundheit zu begrenzen.

Bei Infektionskrankheiten ist der Zusammenhang zwischen Exposition und Wirkung regelmäßig sehr eindeutig, wie auch bei allen anderen Krankheiten, bei denen zwischen Exposition und Auswirkung kurze Zeiträume liegen.

Das besondere Interesse des Epidemiologen bezieht sich auf zwei Tatsachen: Die eine betrifft den Zusammenhang zwischen dem Auftreten von Erkrankungen und der Dauer der Exposition. Der andere betrifft die Relevanz dieser Zusammenhänge jeweils für die epidemiologischen Potenziale. Das Zeitintervall, in dem Menschen einer Exposition ausgesetzt sind, hängt sowohl von den Merkmalen der Exposition, der Art der spezifischen Krankheit als auch von den sozialen und individuellen Bedingungen der Betroffenen ab. Dieses Intervall kann jedoch durch konkurrierende Gefahren, durch die Aggregation verschiedener anderer Gefährdungen oder durch präventive Einflüsse, die die Gefährdung zwar nicht beenden aber reduzieren, verringert werden. Mit anderen Worten:

Im Falle übertragbarer Krankheiten wird diese Spanne auch als Inkubationszeit bezeichnet.

Die Zeit unter Risiko ist von der Zweitspanne zwischen dem Auftreten eines ersten Erkrankungsfalls und dem Beginn einer Epidemie, also der epidemiologischen Latenzzeit zu unterscheiden.

3 Die Klassifikation von Epidemien und Regressionen

Klassifikationen von Epidemien und Regressionen ordnen die Vielfalt je nach deren Ursachen.

Die Reproduktionsziffer nach Ross legt den Versuch nahe, Epidemien und Regressionen nach der Vielzahl ihrer ursächlichen Mechanismen zu klassifizieren. Erste Versuche dieser Art gehen auf B. Kreuz (1919-1981) und Mitarbeiter (*B. Kreuz, B. et al: Über den Zusammenhang zwischen Krankheitsdauer und Morbidität. Ztschr. ärztl. Ausbildung* 67 (1973) 144-148) zurück und sind von Niehoff und Li im Jahr 2022 fortgeführt worden (*Niehoff, J.U., Li, B. Epidemics and Regressions. Saluscon Akademie Verlag, Berlin 2022, ISBN 978-3-948267-15-5*). Unser Klassifikationsvorschlag folgt einem mehrachsigen Konzept, dessen Variablen die Inzidenz, die Falldauer, die Intensität einer Exposition, die Vulnerabilität der jeweiligen epidemiologischen Potenziale, die Durchdringung einer Bevölkerung einer Exposition und die Zeit unter Risiko sind. Unser Modell kennt dann jeweils die Konstanz einer Variablen, ihre Zunahme, ihre Abnahme sowie die vorübergehende Zu- bzw. Abnahme.

Klasse	Variablen						
	Prävalenz	Inzidenz	Falldauer	Intensität der Exposition	Anfälligkeit des Potenzials	Durchdringung der Bevölkerung	Zeit unter Risiko
Epidemie Typ 1	↗	↗	=	↗	=	=	=
Epidemie Typ 2	↗	↗	=	=	↗	=	=
Epidemie Typ 3	↗	↗	=	=	=	↗	=
Epidemie Typ 4	↗	↗	=	=	=	=	↗
Diagnostische Epidemie Typ 1	↗	↘	↗	=	=	=	=
Diagnostische Epidemie Typ 2	↗	→	=	=	=	=	=
Therapeutische Epidemie Typ 1	↗	=	↗	=	=	=	=
Therapeutische Epidemie Typ 2	↗	=	→	=	=	=	=

Abb. 3: Klassifikation von Epidemien

In dieser Tabelle werden die verschiedenen Möglichkeiten für einen Anstieg der Prävalenz klassifiziert. Sie zeigt, dass es mehr Gründe für einen Anstieg der Prävalenz gibt als nur den Anstieg der Gefährdungen. Die unterschiedenen Typen von Epidemien zu veranschaulichen wir wie folgt:

Epidemie Typ 1	Diese Art der epidemischen Exposition wird durch eine Zunahme der Intensität einer bekannten Exposition verursacht, weil sich ihre Merkmale ändern. Diese Veränderungen können zu einer steigenden Inzidenz führen, während alle anderen Auswirkungen gleichbleiben. Als Beispiel kann eine Senkung der Schutznormen für Arbeitnehmer in gefährlichen Industriezweigen herangezogen werden.
Epidemie Typ 2	Diese Art von Epidemie wird durch einen Anstieg der Anfälligkeit der Bevölkerung für die Exposition verursacht. Diese Auswirkungen führen zu einem Anstieg der Inzidenz aufgrund einer Veränderung des Gesundheitszustands der exponierten Bevölkerung, während alle anderen Auswirkungen gleichbleiben. Als Beispiel kann eine Verschlechterung der Impfrate oder des durchschnittlichen Gesundheitszustands dienen.
Epidemie Typ 3	Diese Art von Epidemie wird durch eine Zunahme der Verbreitung der Exposition in einer Bevölkerung verursacht. Dies kann auf eine Verschlechterung der Präventionsmaßnahmen oder auf die Ausweitung von Kontakten mit Exposition zurückzuführen sein, z. B. im Falle übertragbarer Krankheiten. Als Beispiel kann die zunehmende Reisetätigkeit angeführt werden.
Epidemie Typ 4	Diese Art von Epidemie wird durch eine Verlängerung des Zeitraums verursacht, in dem eine bestimmte Bevölkerung gefährdet ist. Die Verlängerung des Zeitraums, in dem die Menschen gefährdet sind. So können selbst bei geringer Exposition die Inzidenzen steigen, auch wenn die Gefahren gering sind.
Diagnostische Epidemie Typ 1	Diese Art von Epidemie wird durch einen Anstieg der Inzidenz aufgrund von Screening-Programmen verursacht, die auf eine frühere Erkennung von behandelbaren Krankheiten oder präventiven Zielen abzielen. Dies hat zeitliche Auswirkungen oder verkürzt das durchschnittliche Diagnoseintervall in einer Bevölkerung. Die Inzidenz steigt vorübergehend an. Als Beispiel kann ein Früherkennungsscreening dienen.
Diagnostische Epidemie Typ 2	Diese Art von Epidemie wird durch einen Anstieg der Inzidenz verursacht, der auf veränderte Normen zurückzuführen ist, die den Status von "gesund" und "krank" unterscheiden. Dies führt zu einem Anstieg der Fälle, die schließlich auf einem neuen Niveau bleiben. Präventive Ziele können als Beispiel dafür gesehen werden.
Therapeutische Epidemie Typ 1	Diese Art von Epidemie wird durch eine Verlängerung des therapeutischen Intervalls verursacht. Dies kann die Folge von Therapien sein, die das Leben verlängern, aber nicht zur Heilung führen. Eine häufig verwendete Messmethode hierfür sind die potenziellen verlorenen Lebensjahre (PYLL), die in einigen Ländern als Standardanbieter für den Einsatz neuer Therapiemethoden verwendet werden.
Therapeutische Epidemie Typ 2	Diese Art von Epidemie bezieht sich auf den verbesserten Zugang und die Nutzung bereits etablierter therapeutischer Methoden. Mehr Patienten werden von diesem Zugang profitieren, was sich auf die Verteilung des therapeutischen Intervalls innerhalb der Gesamtbevölkerung auswirkt. Der gleiche Effekt ist möglich, wenn der Zugang zur Behandlung aufgrund einer sozialen Krise verweigert wird.

Tab. 4: Klassifikation von Epidemien

Regressionen sind die Spiegelbilder von Epidemien.

Klasse	Variablen						
	Prävalenz	Inzidenz	Falldauer	Intensität der Exposition	Anfälligkeit des Potenzials	Durchdringung der Bevölkerung	Zeit unter Risiko
Regression Typ 1							
Regression Typ 2							
Regression Typ 3			 				
Regression Typ 4							
Diagnostische Regression Typ 1							
Diagnostische Regression Typ 2							
Therapeutische Regression Typ 1							
Therapeutische Regression Typ 2							

Abb. 4: Klassifikation von Regressionen

Nachfolgend werden die Möglichkeiten für eine Verringerung der Prävalenz klassifiziert. Sie Übersicht zeigt, dass es mehr Gründe für eine Abnahme der Prävalenz gibt als nur die der Gefährdungen.

Regression Typ 1	Diese Art der Regression wird durch eine abnehmende Intensität der Exposition verursacht. Die Änderung von Produktbestandteilen oder das Verbot von Pestiziden in der Landwirtschaft kann ein Beispiel dafür sein.
Regression Typ 2	Diese Art der Regression wird durch eine abnehmende Anfälligkeit der Bevölkerung für die Exposition verursacht. Ein Beispiel hierfür ist die Verbesserung der Impfrate oder ein besserer allgemeiner Gesundheitszustand der Zielbevölkerung.
Regression Typ 3	Diese Art der Regression wird durch eine Verringerung des Eindringens einer Exposition in eine Bevölkerung oder durch die Begrenzung der Kontakte mit der

	Exposition verursacht, z. B. im Falle der Begrenzung der Exposition durch Sanierung der Lebensbedingungen.
Regression Typ 4	Diese Art der Regression wird durch eine Verringerung der Zeit, in der eine Population gefährdet ist, verursacht. Dies kann auf die Ausweitung der Schutzvorschriften oder auf veränderte biologische (z. B. Alter) oder soziale Merkmale einer Population zurückzuführen sein.
Diagnostische Regression Typ 1	Diese Art der Regression wird durch einen Rückgang der Inzidenz aufgrund einer Verlängerung der Diagnoseintervalle verursacht. Dies kann der Fall sein, wenn sich die Zugänglichkeit der medizinischen Versorgung verschlechtert.
Diagnostische Regression Typ 2	Diese Art der Regression wird durch einen Rückgang der Inzidenz aufgrund von Änderungen der wissenschaftlichen Normen verursacht, die den Status von "gesund" und "krank" unterscheiden. Dies führt zu einem Rückgang der diagnostizierten Fälle pro Zeiteinheit.
Therapeutische Regression Typ 1	Diese Art der Regression wird durch eine Verkürzung des Behandlungsintervalls verursacht. Dies kann der Fall sein, wenn die Behandlungsdauer aufgrund eines schlechten Zugangs zur medizinischen Versorgung infolge einer sozialen Krise verkürzt wird.
Therapeutische Regression Typ 2	Diese Art der Regression wird z. B. durch eine Verkürzung des therapeutischen Intervalls aufgrund besserer Therapiemöglichkeiten verursacht.

Tab. 5: Klassifikation von Regressionen

In der epidemiologischen „Realität“ wird keine dieser idealtypischen und hier zu Klassen geordneten Phänomene isoliert zu beobachten sein. Vielmehr ist davon auszugehen, dass die diesen Phänomenen zugeordneten Ursachen nebeneinander, nacheinander und mit unterschiedlichen Intensitäten sowohl innerhalb eines jeden identifizierbaren epidemiologischen Potenzials wie auch zwischen diesen wirksam sind. Das bewirkt vielfache Differenzierungen der Verläufe von Epidemien und Regressionen. Das Resultat sind dann formal und ursächlich ungleiche Verläufe, die sich dem Beobachter darstellen. Sie werden dann auch bei der Bewertung und der Ableitung von Konsequenzen zu Problemen und Widersprüchen führen.

4 Quantifizierung in der Epidemiologie

Die empirische Grundlage epidemiologischer Forschung sind Daten über gesundheitsrelevante Ereignisse, Zustände und jeweils deren Dauer.

Entsprechende Messungen dienen der Ermittlung absoluter und relativer Häufigkeiten und der Schätzung von Ereigniswahrscheinlichkeiten. Sie dienen der Schätzung von aktuellen und erwartbaren Bedarfen, räumlichen und zeitlichen Vergleichen von Ereignissen und von Wahrscheinlichkeiten für das Eintreten von Ereignissen. Die für die Epidemiologie relevanten Daten bilden dann Querschnitte, Längsschnitte und Zeitreihen ab und können allgemein in Anteilsziffern, Verhältnisziffern, Wahrscheinlichkeiten und Intensitätsmaße unterschieden werden. Die Daten stammen aus vorliegenden Dokumentationen (Sekundärdaten) oder aus Studien (Primärdaten). Soweit solche Daten regelmäßig und auf der Basis hoheitlicher Interessen erheben werden, handelt es sich um ein epidemiologisches Surveillance.

Epidemiologische Forschung beginnt sinnvollerweise mit einer Frage, oder einer Hypothese. Vorliegende Datensammlungen können Fragen veranlassen und Hypothesen begründen. Wenn das Interesse über Ressourcenallokationen hinausreicht, bedürfen Daten logisch sinnvoller Bezüge je entsprechend den verfolgten analytischen Interessen.

Grundsätzlich beschreiben epidemiologische Daten die Verteilung gesundheitsrelevanter Ereignisse, Zustände und Dauern in Bevölkerungen und deren Gliederungen. Folgerichtig haben demografische und epidemiologische Studien auch viele methodische Gemeinsamkeiten (siehe Heft 2). Die analytische Methodik gründet vor allem auf den Methodiken des Vergleiches. Dieser setzt die Definition der Zielgröße (Nominator), deren Zählung/Messung sowie dann der Zuordnung zu ihrem logisch zugehörigen Denominator voraus. Dieser ist dann das räumliche, zeitliche und strukturell (biologisch und sozial) zugehörige epidemiologische Potenzial. Dieses entspricht den von Expositionen gegeneinander abgrenzbaren und von definierten Gefährdungen oder Gefährdungsfolgen Betroffenen. Neben der Intensität einer Exposition sind dann die unterscheidbaren (klassifizierbaren) Dauern, die von einem Potenzial gemeinsam erlebt werden, wichtige epidemiologische Sachverhalte. In der epidemiologischen Forschung ist die sachgerechte Charakterisierung des Denominators oft schwieriger als die des Nominators. Sie ist typischer Weise das zentrale Analyseziel.

Beispiel:

In der Vergangenheit enthielten Bremsbeläge von Kraftfahrzeugen Asbest. Zu diesen Zeiten wurde der Verkehr in Großstädten noch direkt durch Polizisten geregelt. Es dürfte kaum realistisch sein, die „Zeit unter Risiko“ bzw. Exposition retrospektiv auch nur annähernd korrekt zu ermitteln. Da durch Asbestbelastung verursachte

Mesotheliome eher regelhaft zur Inanspruchnahme ärztlicher Leistungen führen und auch richtig diagnostiziert werden, ist deren Zählung ein zu beherrschendes Problem (Nominator). Sehr viel schwieriger, ggf. unmöglich ist hingegen die Zählung der zugehörigen Risikobevölkerung, die einer Asbestbelastung ausgesetzt war und dann auch die Messung der Expositionsdauer (Denominator) des Potentials der Gefährdeten.

In der Epidemiologie sind nur drei Sachverhalte messbar: Ereignisse, Zustände und die Zeit (Dauer).

Ereignisse führen zur Änderung von Zuständen und werden kumuliert als Ereignismengen für einen Zeitraum gezählt.

Zustände beschreiben die stabilen Eigenschaften eines zu messenden Sachverhalts zwischen zwei Ereignissen. Ihre Zählung erfolgt für Zeitpunkte.

Dauern beschreiben die Intervalle zwischen zwei Ereignissen und damit die Dauern von Zuständen.

Ein krankhafter Prozess lässt sich regelhaft durch die Abfolge von Teilereignissen beschreiben, die jeweils in qualitativ neue Zustände führen, bzw. übergehen. Zwischen diesen Ereignissen liegen dann Intervalle mit einer messbaren Dauer.

Beispiel	Ereignis	Intervall/Dauer
Krankheit	Feststellung eines Tumors	Überlebenszeit nach Diagnosestellung
Gesundheitsgefahren	Beginn des Rauchens	Dauer zwischen Expositionsbeginn und Diagnosestellung einer Folgekrankheit
Infektion	Kontakt mit einem Virus	Inkubationszeit, Dauer der Infektiosität
Therapie	Beginn einer Insulintherapie	Überlebenszeit nach Therapiebeginn
Epidemie	Beginn der Zunahme der Prävalenz	Zeitspanne zwischen einem ersten Fall und dem Beginn einer Eskalation der Fallmengen (epidemiologische Latenz)

Regression	Beginn der Abnahme der Prävalenz	Zeitspanne zwischen dem Maximum der prävalenten Fallmenge und dem Übergang in eine neuerliche epidemiologische Latenz.
------------	----------------------------------	--

Tab. 6: Beispiele für epidemiologisch relevante Ereignisse, Zustände und Intervalle

Der Zeitpunkt des Ereignisses "Krankwerden" bzw. die Dauer des Zustands "Kranksein" lassen sich in der Regel nicht exakt bestimmen. Zählbar sind hingegen die Ereignisse der Diagnosestellung, sofern die Krankheit zur Inanspruchnahme medizinischer Leistungen führt und das Ende einer Therapie. Voraussetzung ist die Dokumentation solcher Daten und die technische und rechtliche Möglichkeit solche Daten zum epidemiologischen Fall zusammenzuführen. (Die weltweiten Bemühungen um eine Digitalisierung medizinischer Versorgungsdaten wird hierfür neue Möglichkeiten schaffen; siehe Heft 10). Folglich sind neu diagnostizierte Fälle, sieht man von Ausnahmen ab, immer Fälle, die bereits prävalent sind. In der Dokumentation erscheinen sie formal zwar als inzidente Fälle, tatsächlich wären sie aber logisch richtig der Prävalenz zuzurechnen. Die jeweils dem Krankwerden und dem Ausscheiden aus der Prävalenz zuzuordnenden Zeitpunkte sind folglich nur Schätzgrößen für den Beginn und für das Ende einer Krankheit. Die Güte solcher Schätzungen wird bei der Vielzahl der Krankheiten immer verschieden sein.

Ereignisse werden in der Regel für Zeiträume kumuliert und Zustände als Durchschnitte für die Zeitpunkte innerhalb eines Zeitraums angegeben.

Die Zeitmessung ist für den Epidemiologen häufig eine methodische Herausforderung. Es kann zumeist nicht ermittelt werden, wie viele Personen exakt im Zeitraum t_{1-0} krank geworden sind oder über gleiche Zeiträume hinweg und im Altersbezug Gefährdungen ausgesetzt waren. Dies ist realistischer Weise auch nicht für die Menge der Personen möglich, die im Zeitpunkt t_x krank sind. Korrekt bezeichnet sind deshalb die sog. Neuerkrankungen, in der Regel die neu diagnostizierten Kranken, bzw. sind sie ein Anteil aller bereits dokumentierten Fälle. Es wird also nicht das Ereignis „Krankwerden“ gezählt, sondern das Ereignis „Diagnosestellung“ zu einem Zeitpunkt, indem die epidemiologischen Fälle bereits der Prävalenz zugehörig sind. Dabei sind immer auch die Mengen der falsch positiven und der falsch negativen Fälle abzuschätzen und im Falle von Stichproben, die Konfidenzintervalle zu ermitteln. Dies ist erforderlich, weil die gewonnenen Daten der Stichprobe in ihrem Verhältnis zur Wirklichkeit immer nur innerhalb der Konfidenzintervalle als richtig zu akzeptieren sind. Hieraus folgt, dass Zufallsdaten, die nicht aus repräsentativen Stichproben resultieren in ihrer Konfidenz nicht beurteilt werden können. Sie können folglich in epidemiologischen Kontexten auch nicht bewertet werden.

Die Angabe neu diagnostizierter Erkrankungen erfolgt für Zeiträume, üblicherweise für ein Kalenderjahr oder für Teile eines Kalenderjahres. In solchen Fällen ist dennoch der korrekte Bezug zur kalendarischen Zeit notwendig. Aus praktischen Gründen werden für solche Angaben reale oder fiktive Zeitpunkte vereinbart. Solche realen Zeitpunkte sind z.B. der Beginn oder das Ende eines Beobachtungszeitraums, bspw. eines Quartals oder Kalenderjahres. Ein fiktiver Zeitpunkt ist dann die mittlere Anzahl der Fälle im Beobachtungszeitraum unter der Annahme, diese Anzahl entspräche mit hinreichender Genauigkeit auch jeweils der Anzahl zu jedem anderen Zeitpunkt.

Wegen der Heterogenität von Bevölkerungen werden solche Verteilungen in jedem der relevanten epidemiologischen Potenziale verschieden sein, sich also als Unterschiede zwischen diesen darstellen. Entsprechende Unterschiede und deren Veränderungen bedürfen dann ursächlicher Erklärung. Erklärungs- oder interpretationsbedürftig sind sie dann, wenn sie (im statistischen Sinne) nicht zufällig sind. Ist dies der Fall, müssen solche Unterschiede dann auch in ihrer praktischen Relevanz bewertet werden.

4.1 Das Allgemeine epidemiologische Krankheitsmodell

Das Allgemeine epidemiologische Krankheitsmodell dient dem Zweck, den Verlauf von Krankheiten zu modellieren.

Krankheiten können über ihren Verlauf (Pathogenese), also über eine zeitliche Abfolge wechselnder Zustände beschrieben werden. Jeder dieser Zustandswechsel ist formal ein Ereignis. Die „Ereignisse“ und die „Intervalle“ zwischen ihnen sind Gegenstände vielfältiger sowohl deskriptiver wie analytischer Fragestellungen.

In epidemiologischen Analysen, die eine Zielkrankheit immer an vielen Personen/Fällen untersuchen, stellen sich diese Punkte und Intervalle dann in einer Vielzahl von Zustandsformen oder Ausprägungen, bzw. in Verteilungen dar. Diese sind für einzelne Krankheiten dann gleichsam die Spiegelbilder der in einer Bevölkerung ablaufenden krankhaften Prozesse - oder auch der Einflussnahmen auf sie durch Prävention oder durch die medizinische Versorgung. Es werden folglich auch alle relevanten Gegenstände der sozialmedizinischen Versorgungsforschung durch das Modell abgebildet.

Epidemiologische Studien machen es erforderlich, die relevanten Aspekte eines krankhaften Prozesses durch geeignete Variablen und dann jeweils deren Häufigkeit, Ausprägung und Verteilung in der jeweils untersuchten Studienpopulation abzubilden.

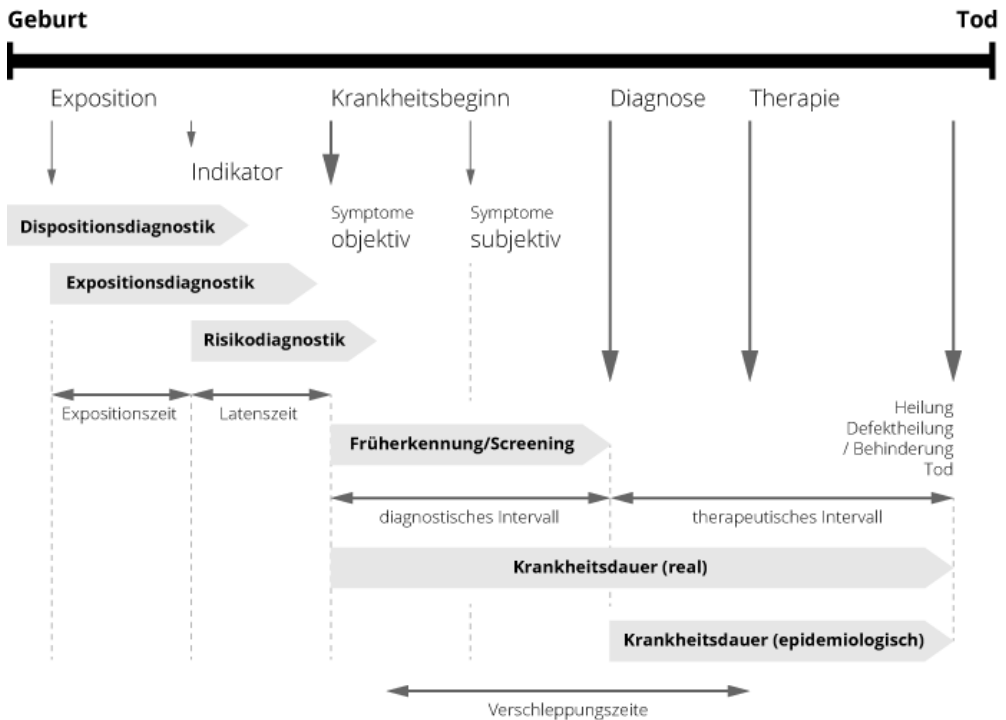


Abb. 5: Das Allgemeine Sozialmedizinische Krankheitsmodell (Niehoff, J.-U. Der Morbiditätsprozess. Habilitationsschrift, Humboldt Universität Berlin 1983)

Das allgemeine Krankheitsmodell bezeichnet alle für die epidemiologische Analytik grundlegenden zeitlichen Zuordnungen von Ereignissen und Zuständen:

1. den Beginn einer Krankheit
2. die Inanspruchnahme von Hilfeleistungen und die Diagnosestellung
3. den Verlauf einer Krankheit unter Therapie
4. das jeweils zurechenbare medizinisch-präventive Intervalle
5. das diagnostische und das therapeutische Intervall
6. die sog. Verschleppungszeiten

Das Modell verdeutlicht die enge sachliche Beziehung epidemiologischer Forschungsansätze zu denen der Versorgungsforschung.

4.2 Zeitpunkte und Intervalle

Für die Epidemiologie ist der Beginn einer ursächlichen Einwirkung (Exposition), deren Dauer, der Beginn einer Krankheit, dann die Diagnosestellung und die Art und der Zeitpunkt des Endes einer Krankheit von analytischer Bedeutung.

Welche Zeitpunkte bzw. Intervalle jeweils im Mittelpunkt des medizinischen Interesses stehen, hängt von der konkreten Fragestellung ab. Interessierende und für theoretische, wie praktische Schlussfolgerungen wichtige Zeitpunkte und Intervalle sind z.B.:

1. der Beginn von Veränderungen aus denen ggf. Krankheit entsteht (z.B. sog. Borderline-Hypertonie, carcinoma in situ)
2. die Dauer des Überganges von diskreten Veränderungen zu klinisch manifester Krankheit (Latenzzeit, bei übertragbaren Krankheiten Inkubationszeit)
3. der Zeitpunkt, zu dem Krankheit erstmalig objektivierbar ist
4. das Intervall bis zur subjektiven Wahrnehmung erster Symptome
5. der erste Kontakt mit dem Arzt
6. Verschleppungszeiten
7. Zeitpunkt der Diagnosestellung
8. Beginn und Ergebnisse der Therapie
9. Ende der Therapie

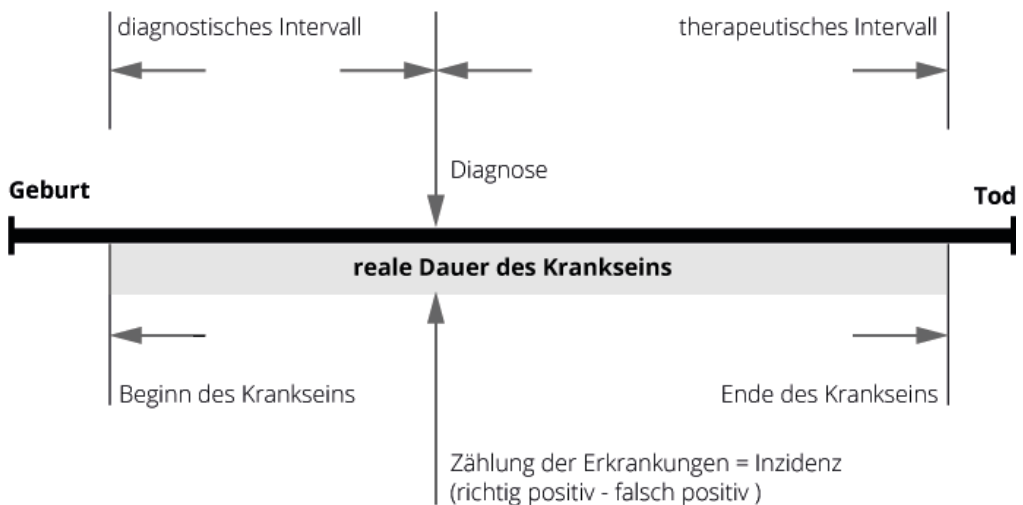


Abb. 6: Diagnostische und therapeutische Intervalle

Die Abbildung zeigt, dass einige wichtige Zeitpunkte und Intervalle innerhalb eines krankhaften Prozesses nicht direkt beobachtbar sind, wie z.B. der Beginn und folglich auch nicht die Dauer einer Krankheit. Es ist zudem wichtig zu verstehen, dass die „Dauer einer Krankheit“ keine Eigenschaft der Krankheit, sondern eine der kranken Personen ist. Diese Intervalle sind also u.a. vom Alter zum Zeitpunkt des Krankwerdens, der Selbstwahrnehmung des Patienten oder der Zugänglichkeit und Qualität des Versorgungssystems abhängig. Es ist zu beachten, dass die diagnostischen und die therapeutischen Intervalle variabel sind und deren Verschiebung auf der Zeitachse immer auch erhebliche Mengeneffekte haben wird. Aus sozialmedizinischer Sicht ist herauszustellen, dass immer auch von einer sozialspezifischen Verteilung solcher Intervalle auszugehen ist, deren Abbildung bspw. medizinische Versorgungsprobleme offenlegen kann.

Der Nutzer der Epidemiologie sollte immer berücksichtigen, dass Schätzungen der Krankheitsdauer eher dann schlecht sind, wenn

1. eine Krankheit nur selten zu ärztlicher Inanspruchnahme führt
2. die Definition einer Krankheit sich immer wieder ändert
3. das diagnostische Intervall lang ist
4. das Intervall zwischen Expositions- und Krankheitsbeginn lang ist
5. Änderungen in Diagnostik und/oder Therapie diese Intervalle verändern
6. der Beginn der Inanspruchnahme medizinischer Hilfe in einer Bevölkerung z.B. aus sozialen Gründen sehr unterschiedlich verteilt ist

Da für Risikoschätzungen aber der Krankheitsbeginn oder z.B. für Versorgungsstudien die Dauer des Behandlungsbedarfs wichtige Messgrößen sind, werden diese zu meist aus dem Zeitpunkt der Diagnosestellung geschätzt. Der Ausweg können klinische Studiensettings sein, die diese Zeitpunkte unter quasi-experimentellen Bedingungen messen. Dies ermöglicht allerdings nicht die Abbildung bevölkerungsbezogener Verteilungen.

Auch dieser Aspekt unterscheidet epidemiologische und klinische Studien: Die Wirksamkeit einer Intervention unter den Bedingungen einer kontrollierten klinischen Studie (efficacy) und die Wirksamkeit unter den Bedingungen der Versorgungswirklichkeit (community effectiveness) bilden unterschiedliche Realitäten ab.

4.3 Die epidemiologischen Studientypen

4.3.1 Studien mit Primärdaten

Es ist immer anzustreben, dass die Datenerhebung für analytische Studien in der Hand des Untersuchers liegt.

In der epidemiologischen Forschung hat sich eine Reihe von speziellen Studientypen herausgebildet. Diese Studientypen werden nachfolgend nur als Übersicht dargestellt. Für weitergehende Informationen sollte der Nutzer auf spezielle Literatur zurückgreifen.

Die epidemiologischen Studientypen lassen sich nach methodischen Kriterien zu insgesamt 4 Gruppen zusammenfassen, die häufig auch in Kombinationen vorzufinden sind. Diese sind

1. Querschnittstudien (engl. Survey, häufigster Studientyp; für deskriptive und Erkundungsstudien geeignet)
2. Längsschnittstudien (aufwändigster Studientyp mit prospektiver Datengewinnung besonders zur Hypothesentestung geeignet)
3. Fall-Kontrollstudien, bei denen bereits an einer Zielerkrankung Erkrankte mit Personen, die nicht an der Zielkrankheit leiden, verglichen werden (Expositionsstudien)
4. Simulationsstudien und Studien zur Prozessmodellierung (finden bspw. Anwendung in Untersuchungen zur prospektiven Evaluation und Folgeabschätzung von bevölkerungsbasierten Interventionen etwa im Rahmen der Versorgungsfor schung; sie sind auch eine wichtige Methodik in retrospektiven Studien zum Verteilungswandel von Gesundheitsproblemen in umschriebenen Bevölkerungen für die keine oder nur wenige Primärdaten rekonstruiert werden können).

In Totalerfassungen wird eine gesamte Zielbevölkerung in die Studie eingeschlossen (nur für ausgewählte Personengruppen und von Hypothesen geleitete Fragestellungen vertretbar). Häufiger sind Stichprobenstudien (die Ziehung einer repräsentativen Stichprobe verlangt gute Kenntnisse der Variablenverteilung in der Grundgesamtheit).

Analytische Studien sind hypothesengeleitet. Sie zielen auf Sachaufklärung beobachteter Phänomene und auf die Testung von Hypothesen. Sie sind analytisch, weil sie ereignisorientiert Ursache-Wirkungs-Zusammenhänge z.B. mittels Kohortenstudien oder in Fall-Kontrollstudien untersuchen.

Interventionsstudien sind eine Sonderform der Kohortenstudien. Viele Interventionsstudien sind den klinisch-experimentellen Bevölkerungsstudien zuzuordnen. Andere weisen Schnittstellen zu epidemiologischen Studien auf, beispielsweise Untersuchungen zur Wirksamkeit von bestimmten Früherkennungsmaßnahmen oder bevölkerungsbezogenen Präventionsansätzen. Sofern es sich dabei um experimentelle oder quasi-experimentelle Bevölkerungsstudien, bzw. um Studien an großen Bevölkerungsstichproben handelt, können diese auch als interventive Kohortenstudien der interventionellen Epidemiologie aufgefasst werden. Die ethischen Probleme sind oft erheblich, wenn bspw. die individuellen Einverständnisse der Bevölkerung nicht eingeholt werden.

Die Güte einer epidemiologischen Studie lässt sich an folgenden Prüfkriterien messen:

1. Konkordanz von Studienziel und Studiendesign
2. Reliabilität der Messinstrumente
3. Validität und Plausibilität der Ergebnisse
4. Übereinstimmung mit anderen Untersuchungen

Ergebnisse einer Studie sind valide, wenn die Studie exakt den Sachverhalt misst, der gemessen werden sollte. Die Validität des beobachteten Zusammenhanges ist der Wahrscheinlichkeit, dass dieser Zusammenhang durch Zufall, bias und confounding erklärt werden kann, umgekehrt proportional.

Die Reliabilität gibt die Zuverlässigkeit und Präzision einer Messung an, d.h. die Wiederholbarkeit der Ergebnisse am gleichen Studienobjekt, ggf. auch durch andere Untersucher. Die Testung der Reliabilität spielt immer dann eine große Rolle, wenn gesundheitsrelevante und in einfachen Häufigkeitsverteilungen nicht zu erkennende Sachverhalte mit speziell entwickelten Messinstrumenten erhoben werden sollen. Es handelt sich oft um die Messung qualitativer Eigenschaften (z.B. mittels Fragebögen zur Selbstwahrnehmung von Krankheiten), die über Skalierungen, Entscheidungen oder Zuordnungen quantifiziert werden sollen. Immer dann also, wenn gleiche Sachverhalte von Studienteilnehmern oder Untersuchern mittels einer Messskala beschrieben werden, muss mit einer Ergebnisvarianz gerechnet werden, die durch den Messvorgang und seine Instrumente selbst erzeugt wird. Beurteilungen der Reliabilität zielen also auf die Reproduzierbarkeit von Studienergebnissen bei Nutzung jeweils gleicher Messinstrumente.

4.3.2 Studien mit Sekundärdaten

Epidemiologische Studien, die auf bereits vorhandene aber für andere Zielstellungen erhobene Daten zurückgreifen, nutzen diese sekundär.

Die Auswahl der Gegenstände der Datensammlungen erfolgt in der Regel administrativen Dokumentationsbedürfnissen und Vorgaben und zwingt zur sorgfältigen Prüfung der Informationsgehalte.

Ein typisches Anwendungsfeld ist die Gesundheitsberichterstattung.

Die Gesundheitsberichterstattung ist eine selektive, problembezogene und wertende Sammlung von wesentlichen Gesundheitsproblemen und Gesundheitsrisiken einer räumlich, zeitlich und/oder alters-, bildungs-, sozialschichtspezifisch definierten Bevölkerung. Sie dient der Information der Öffentlichkeit und/oder der Politik.

Typische Datenquellen für Sekundärdaten sind:

1. Medizinalstatistiken
2. Todesursachendokumentationen
3. Dokumentationen zur Arbeitsunfähigkeit

4. Dokumentationen meldepflichtiger Krankheiten
5. Dokumentation verordneter Medikamente
6. Leistungsstatistiken von Krankenversicherungen und Erbringern medizinischer Leistungen
7. Dokumentationen der Renten- und Unfallversicherungsträger
8. Krankheits- oder expositionsbezogene Register
9. Befragungs- und Paneldaten

Sekundärdaten werden weltweit genutzt und sind – auch bei berechtigten Zweifeln an ihrer Aussagefähigkeit im einzelnen Fall – oft die einzigen und global auch eher für wenige Bevölkerungen verfügbaren Informationsquellen zu den Gesundheitsproblemen von Bevölkerungen.

4.4 „Big Data“ in der epidemiologischen Forschung

„Big data“ bezeichnet riesige unstrukturierte und oft weltweit generierte und gehandelte Sammlungen von Massendaten aus praktisch allen überwachungs- und dokumentationsfähigen Lebensbereichen. Solche Sammlungen schließen oft Daten zur Gesundheit ein oder lassen auf sie Rückschlüsse zu.

„Big data“ erzeugen aktuell erhebliche Interessen an ihrer kommerziellen und/oder wissenschaftlichen Verwertbarkeit. Diese Datensammlungen sind oft die Nebenprodukte technischer Entwicklungen, die zur Sammlung von Daten führen und dann deren unreflektierte Speicherung veranlassen. Sie sind mittels automatisierter Algorithmen und vom Menschen unabhängiger Klassifikationsprozeduren quasi zeitgleich verarbeitungsfähig. Mittels der Konzepte der künstlichen Intelligenz ist diese zunehmend selbst in der Lage, dann auch Entscheidungen zu treffen. Damit verlässt „big data“ die Ebene der Analytik und wird potenziell zum handelnden Akteur (Medizin 4.0).

Autonom entstehende Datenberge aus solchen Sammlungen werfen die Frage auf, ob sie auch für epidemiologische Fragestellungen nutzbar sein können. Ein solches Verfahren kann auch als „data mining“ in den Mengen von Sekundärdaten aufgefasst werden. Die Gründe ihrer Entstehung, die Kriterien ihrer Dokumentation und die Datenqualitäten sind in der Regel ebenso wenig zu kontrollieren wie ihre Selektivität und Repräsentativität. Die Tauglichkeit solcher Daten für epidemiologische Studien muss deshalb im Einzelfall sorgfältig geprüft werden. Die Eigentümer dieser Datenmengen gehen davon aus, dass es ein weitreichendes wirtschaftliches Interesse an der Abbildung und ggf. Steuerung individueller Verhaltens-, Interessen- und Risikoprofile gibt. Da die Datenquellen und die Auswertungsverfahren häufig dem Eigentumsrecht unterliegen, genügen auch die Resultate eher selten den Standards wissenschaftlicher Studien.

Auch aus der Perspektive sozialmedizinischer und speziell epidemiologischer Analyseinteressen ist das Interesse an Daten groß. Unübersehbar sind allerdings auch die Schwierigkeiten, big data zu bewerten und zu nutzen. Die Datenqualität (z.B. Validität, Reliabilität, Präzision, differentielle und nicht-differentielle Fehlklassifikationen, ggf. Sensitivität, Spezifität, positive und negative prädiktive Werte, Selektivität, Repräsentativität etc.) dürfte kaum beurteilbar sein, da es nicht realistisch ist, anzunehmen, es ließen sich globale Standards für die Datenbeschreibung, die Klassifikationen der Daten und vor allem für die Beschreibung der zugehörigen Grundgesamtheiten vereinbaren.

Es ist anzunehmen, dass diese Datenkonvolute dennoch eine erhebliche praktische Bedeutung erlangt haben und zunehmend erlangen werden. Grundlegende Probleme sind jedoch bezüglich mindestens dreier Schlüsselanforderungen zu benennen: diese sind die unsichere und in aller Regel geringe Qualität der Daten, die Unklarheit des Bezuges der Daten zu ihrer Grundgesamtheit und die praktischen Konsequenzen der hohen Erneuerungsgeschwindigkeit der Daten. Unbezweifelbar bilden sie die Realitäten ab, doch kann oft kaum geklärt werden, welche „Realitäten“ durch sie abgebildet oder nur erzeugt werden.

Beispiel:

Es ist technologisch bereits heute und dies mit wachsender Geschwindigkeit (Vermutungen sprechen von einem exponentiellen Wachstum) möglich, die medizinischen Versorgungsprozesse einer Region mit hunderttausenden Einwohnern informationsseitig bis auf die Ebene des Alltagslebens von Individuen abzubilden, diese Daten dann zu zentralisieren und Versorgungsprozesse wie das therapiegeeignete Verhalten der Individuen wie auch die ärztlichen Entscheidungen zu überwachen und zu steuern. Dort wo dies möglich ist und auch heute schon praktiziert wird, stehen Wissenschaftler vor der Frage, ob und ggf. für welche Fragen die gesammelten Überwachungsdaten analytisch nutzbar gemacht werden könnten. Es ist offensichtlich, dass hier nicht allein wissenschaftliche Fragen zur Diskussion stehen. Gemeinsames Merkmal der bekannten Beispiele ist ihr geringer Anspruch an den Schutz der von durch solche Daten berührten Individualrechte.

Die Nutzung von Massendaten ist in der Epidemiologie nicht neu. Neu sind jedoch die Datenmengen, ihre technische Generierung und Verwaltung, deren privatwirtschaftliche, wie auch öffentlich-rechtliche Monopolisierung auch ohne öffentliche Kontrollmöglichkeiten, ihr individueller Bezug bei weitgehendem Verlust des Schutzes von Persönlichkeitsrechten, die Globalität der Sammlungen, ihr globales Trading und die fehlende Regulierung der Datennutzung. Es handelt sich um eine technische Entwicklung, die sich mit den entstehenden Möglichkeiten der künstlichen Intelligenz verbindet und auch für die epidemiologische Forschung, vor allem im Interesse risikoäquivalenter Krankenversicherungen große Bedeutung hat.

Viele dieser Daten werden heute und künftig nicht zur Beantwortung einer wissenschaftlichen Fragestellung gezielt gesammelt. Sie entstehen einfach durch die Nutzung technischer Geräte, z.T. als Verhaltensprotokolle der Internet- und Cloudnutzer aber auch durch den freiwilligen und bewussten Verzicht auf die Eigentümerschaft individueller Daten durch die Nutzer technischer Geräte. Die Erwartung ist, dass solche nicht an Fragen gebundenen also auch kaum strukturierten Daten allein durch ihre Menge und Totalität in der Lage sein könnten, Wissen zu generieren. Diese Erwartung wird von der Vorstellung gespeist, die Verfügbarkeit aller überhaupt nur messbaren Sachverhalte könne dann auch alle Informationen über die Wirklichkeit zumindest potenziell beinhalten. Wer also alle Buchstaben kennt, könne - die technischen Möglichkeiten vorausgesetzt - mit Hilfe der künstlichen Intelligenz auch alle Bücher schreiben, eben durch die endlose Generierung der zufälligen Kombination der Schriftzeichen.

Wichtig ist für den Nutzer zunächst, die Prinzipien wissenschaftlichen Arbeitens nicht zu verletzen und u.a. die Qualität dieser sekundären Massendaten zu testen und zu bewerten. Diese Situation ist nicht neu. Für die grundlegenden Ansprüche an das wissenschaftliche Arbeiten ändert sich durch „big data“ nichts, für die wissenschaftlichen Arbeitsprozesse jedoch Vieles.

Es bleibt das Wesen von Wissenschaft, dass ohne Fragen und Hypothesen kein Wissen geschaffen werden kann. Es gibt also nur Antworten, wenn es zuvor Fragen gegeben hat. Die Entwicklung wissenschaftlicher Fragen ist aber vom wissenschaftlichen Arbeitsprozess selbst nicht zu trennen. Auch unter den Gegebenheiten von „big data“ sind Prinzipien wie die Bewertung der Datenqualität, ihrer Repräsentativität, ggf. deren Randomisierung und Testung auf Validität und Reliabilität einzuhalten.

Beispiel:

Der Umfang und die Qualität einer Bibliothek sind bedeutungslos, wenn sie nie geöffnet ist oder ihre Nutzer des Lesens unkundig sind. Eine solche Bibliothek wäre auch sinnlos, wenn sie allein die von den Völkern der Erde jemals benutzten Schriftzeichen sammeln würde.

„Big data“ ist kein neues Phänomen. Es ist auch nicht neu, dass die Generierung von Fragen mit der von Daten nicht Schritt hält und sich die Diskussion um die Nutzung der Daten stattdessen auf technische Probleme fokussiert.

Neu sind allerdings die Datenvolumina, die zu umfangreich und zu komplex sind, um sie mit den (heute noch) üblichen Techniken der Datenaufbereitung und -analyse (etwa Lesen und Verstehen) zu bewältigen. Diese Datenmengen unterliegen zudem einer ständigen nicht selektierenden, nicht thematisierten, also weitgehend beliebigen und in ihrer Qualität nicht prüfbaren Ausweitung. Wer diese Daten bearbeiten kann, ist der Monopolist des Wissens. Damit gehen alle demokratischen

Rechtskonstruktionen verloren, weil sich die Datenverfügbarkeit und -nutzung der öffentlichen Kontrolle entzieht.

Es ist eine häufige Erfahrung von Wissenschaftlern, gefragt zu werden, was mit existierenden sekundären Datenbeständen zusätzlich Sinnvolles anzufangen sei. Folglich geht dann die Generierung von Daten einem Wissenschaftszweck, also etwa der Formulierung einer Frage oder einer Hypothese voraus. Es ist auch eine Erfahrung, dass die Entstehung von Datenmengen Impulsgeber für technische Entwicklungen zu ihrer Aufbereitung ist.

Beispiel:

Der französische Sozialhygieniker L. R. Villermé (1782-1863) untersuchte selbst 760.000 Arbeiter, um deren Gesundheitsverhältnisse im Kontext ihrer Lebensbedingungen zu beschreiben. In dieser frühen epidemiologischen Studie entstanden Massendaten, für die es zu jener Zeit keine Analysekonzepte gab. Dies war für L. A. J. Quetelet (1796-1874) der Anlass, sich mit methodischen Konzepten für solche Analysen auseinanderzusetzen. Aus heutiger Sicht ist u.a. interessant, dass die Grenzen dieser Versuche zwar auch im methodischen Bereich lagen, vor allem aber dann offensichtlich wurden, wenn Hypothesen und Interpretationen der Daten an die Grenzen des Kenntnisstandes der Zeit stießen.

Die Vermehrung von Datenmengen erzeugt allein also kein neues Wissen. Die Mehrung von Wissen verlangt vielmehr die Qualifizierung von Fragen.

Für die Zukunft ist deshalb nicht „big data“ zu problematisieren, sondern die, dann aber aus anderen Quellen gespeisten, Entwicklungen zur und um die sog. künstliche Intelligenz und die Chancen und öffentlichen Möglichkeiten, deren Anwendung in den verschiedenen Bereichen zu kontrollieren. Hier ist sicherlich auch für die Epidemiologie Vieles zu erwarten.

An der Notwendigkeit einer sehr kritischen Bewertung sekundärer Daten und den daraus abgeleiteten Informationen ist also auch dann nichts zu ändern, wenn diese „big“ sind.

4.5 Ziffern in der Epidemiologie

Die quantitative Beschreibung eines zu untersuchenden Sachverhalts ist die Grundlage jeder epidemiologischen Studie (measurement of occurrence).

In der Epidemiologie heißt messen absolute Häufigkeiten von Ereignissen und Zuständen zu zählen, Dauern zu bestimmen und die Ergebnisse zu deren

epidemiologischen Potenzialen in Beziehung zu setzen. Dies hat mindestens zwei Voraussetzungen:

1. Der Gegenstand der Untersuchung (Ereignis, Zustand oder Dauer) muss exakt bestimmt werden. Hierbei handelt es sich im statistischen Sinne um die Definition des Nominators.
2. Der logische Bezug des Untersuchungsgegenstandes muss benannt werden (Umfang des logisch zugehörigen epidemiologischen Potenzials). Hierbei handelt es sich im statistischen Sinne um die Definition des Denominators.

4.5.1 Die Quantifizierung von Zuständen

Gleich bezeichnete Zustandsmengen entsprechen in der Epidemiologie der Prävalenz.

Die Menge der prävalenten Fälle ist das quantitative und strukturelle Potenzial für die weitere epidemiologische Dynamik der Prävalenz. In der Epidemiologie sind mindestens drei Arten von Prävalenzraten gebräuchlich:

$$\text{Punktprävalenzrate} = \frac{\text{Gefährdete oder Kranke in } t_x}{\text{das epidemiologische Potenzial in } t_x}$$

$$\text{Periodenprävalenzrate} = \frac{\text{Gefährdete oder Kranke in } t_1 - t_0}{\text{mittleres epidemiologisches Potenzial in } t_1 - t_0}$$

$$\begin{aligned} \text{Lebenszeitprävalenzrate (Life - time prevalence rate, (LTP))} \\ = \frac{\text{Gefährdete oder Kranke einer Sterbetafelkohorte bis zum Alter } e_x}{\text{Sterbetafelbevölkerung bis zum Alter } x} \end{aligned}$$

Im Falle übertragbarer Krankheiten sind für den Nominator ggf. zusätzlich die infizierten von den kranken Personen zu unterscheiden.

Die Dokumentation der absoluten Zahlen prävalenter Fälle ist für die Zuweisung von Ressourcen notwendig und auch ausreichend. Bei Vergleichen müssen diese absoluten Zahlen jedoch immer ins Verhältnis zum Nenner gesetzt werden, der die Zielpopulation der Vergleiche darstellt. Der gewählte Nenner ist das epidemiologische Potenzial, aus dem die gezählten Personen stammen. Dies macht es notwendig, zwischen Personen und Fällen zu unterscheiden, wenn einer Person innerhalb eines Beobachtungszeitraums mehr als ein Fall zugeordnet werden kann. Häufig wird die Gesamtbevölkerung als Nenner genommen, da das wahre epidemiologische Potenzial unbekannt bleibt. Es wird also der arithmetische Mittelwert für die Gesamtbevölkerung berechnet.

Prävalenzraten messen den Anteil eines epidemiologischen Potenzials der entweder exponiert, infiziert oder krank ist. In der Zeitreihe dokumentieren diese Raten die quantitative Dynamik der krankheitsspezifischen Prävalenz. Anteilsziffern sind keine Maße für die Gefährdung durch eine Exposition. Sie sind folglich auch keine Maße für die Veränderung von Gefährdungen. Anteilsziffern lassen sich in der Zeitreihe nicht addieren und nur vergleichen, wenn die Strukturgleichheit der Potenziale gesichert ist.

Es muss immer wieder darauf hingewiesen werden, dass Prävalenzraten keine Wahrscheinlichkeiten sind, keine Auskunft über Gefährdungsänderungen geben und auch nicht für Vorhersagen oder für Gefährdungsanalysen verwendet werden können. Um solche Informationen zu erhalten, müssten die Gründe geklärt werden, warum eine Prävalenzrate im Laufe der Zeit zu- oder abnimmt. Zu diesem Zweck müssten also die Ursachen der Dynamik benannt und quantifiziert werden. Hierzu bedarf es der Kenntnis/der Schätzung der Inzidenzen und der Zeitintervalle, in denen die epidemiologischen Fälle den Prävalenzen zugerechnet werden können. Nur wenn das Intervall, in dem die Fälle auftreten, sehr kurz ist, ist es möglich, die Prävalenz als groben Schätzer für die Inzidenz zu verwenden.

Punktprevalenzraten sind nur unter bestimmten Bedingungen zu messen. Das ist möglich, wenn tägliche Meldungen der Prävalenzdynamik gesetzlich oder versicherungsrechtlich vorgeschrieben sind und zu diesen Zeitpunkten auch die zugehörigen Potenziale (Denominatoren) bekannt sind.

Beispiel:

Nach einem vom Robert-Koch-Institut durchgeführten telefonischen Gesundheits-survey 2003 sagten zum Zeitpunkt der Befragung 5,8 %, sie litten an einer Angina pectoris. Einen Herzinfarkt hatten 2,7 % erlitten und eine Arthrose hätten 18,7 %. 61,8 % der Befragten klagten über Rückenschmerzen. Zum Zeitpunkt der Befragung waren 32,5 % Raucher, 13,1 % waren schwerbehindert oder in ihrer Erwerbsfähigkeit gemindert, 22,5 % hatten sich gegen Grippe impfen lassen. Alles dies sind Prävalenzangaben, sofern die Befragung an einem Tag erfolgte, handelt es sich auch um eine Punktprevalenzrate.

Periodenprävalenzraten sind vor allem bei länger andauernden Epidemien von Interesse. Diese Raten geben Aufschluss über die Auswirkungen von Epidemien und Regressionen auf die Bevölkerung und ihre medizinischen Versorgungssysteme. Sie beschreiben folglich die Relevanz eines Problems.

Die Lebenszeitprävalenzrate (LTP) ist eine reale oder alternativ, eine fiktive Kohortenmessung. Sie ist nur mittels spezifisch angepassten Dekrementtafeln (Sterbetafeln) zu berechnen. Diese Rate kann auch als eine auf die Lebensdauer der Potenziale standardisierte Prävalenzrate interpretiert werden. Die LTP setzt voraus, dass die Mengen der erkrankten Personen und die der Fälle gleich groß sind. Dies ist der Fall, wenn die

betreffenden Krankheiten im Lebenszeitraum nicht wiederholbar sind. LTP spielt in einigen spezifischen Krankenversicherungsmodellen eine Rolle.

4.5.2 Die Quantifizierung von Ereignissen

Gleichbezeichnete Fälle des Ereignisses „Erkrankung“ entsprechen in der Epidemiologie der Inzidenz.

Die inzidenten Fälle bezeichnen die Mengen von Zustandsänderungen, also den Übergang von einem in ein anderes Potential. Sie sind in diesem Sinne Bewegungsmengen und eine Ursache für die weitere Dynamik der Prävalenz.

Die dokumentierten Mengen werden als Anteile an der Referenzbevölkerung bzw. als Anteile an den epidemiologischen Potenzialen relativiert, deren Gesamtsumme oder Obermenge die Gesamtbevölkerung ist. Die neuen Fälle treten auf, wenn gesunde Personen exponiert werden oder, wenn exponierte Personen erkranken. Im Falle von Infektionskrankheiten sind das dann Personen, die sich infizieren, bzw. infizierte Personen, die erkranken. Damit sie als Fälle gezählt werden können, müssen sie jedoch zunächst entdeckt bzw. diagnostiziert werden. Es ist immer sicherzustellen, dass die Personen- und die krankheitsspezifischen Fallmengen den Personenmengen, die betroffen sind, entsprechen.

Die Menge der neu dokumentierten Fälle hängt davon ab, ob und wie viele Menschen sich entscheiden, einen Arzt aufzusuchen oder ob es Public Health Aktivitäten gibt, nach ihnen aktiv zu suchen. Es ist deshalb zu prüfen, in welchem Verhältnis die quantitativen Abbilder der Realität zu dieser stehen. Dies ist insbesondere dann schwierig, wenn die Mengen durch soziale Merkmale der Patienten oder durch sozial unterschiedliche Zugänge und Zugangsbereitschaften zur Diagnostik verzerrt sind.

4.5.3 Die Inzidenz

$$\text{Inzidenzrate (I)} = \frac{\text{neu Exponierte, Infizierte oder Erkrankte } t_1 - t_0}{\text{mittlere Anzahl des jeweiligen Potenzials } t_1 - t_0}$$

Diese Rate schätzt den Anteil der neu entdeckten exponierten, infizierten oder erkrankten Personen/Fälle im Verhältnis zu ihrem logisch zugehörigen Potenzial, z. B. zur durchschnittlichen Bevölkerung im Beobachtungszeitraum. Ein Schlüsselproblem für Epidemiologen bleibt aber immer die Frage, ob der Nenner dieser Gleichung in einem logischen Verhältnis zum Zähler steht.

Beispiel:

Es wird von 20 Tuberkulosefällen pro 100.000 Personen in einer Bevölkerung berichtet. Unter der Annahme, dass die Diagnose korrekt ist, bleibt immer noch unklar, ob diese 20 logisch richtig auf 100.000 Personen der Gesamtbevölkerung bezogen werden können. Es könnte sein, dass nur eine einzige und sehr spezielle soziale Gruppe von Menschen, die unter außergewöhnlichen Bedingungen leben, das wahre epidemiologische Potenzial darstellt. Diese Gruppe kann nur 1.000 exponierte Personen umfassen. Es macht dann einen Unterschied, ob das Risikomanagement sich auch 100.000 oder nur auf 1.000 Personen beziehen soll.

Die Inzidenzrate misst also den Anteil der Personen, die in Bezug auf ein definiertes Potenzial erkranken, in dem Beispiel an Tuberkulose. Sie misst keine Wahrscheinlichkeit, an Tuberkulose zu erkranken. Die Rate hängt vielmehr vom Umfang und den besonderen Strukturen der gewählten Referenzpopulation sowie von den diagnostischen Aktivitäten ab. Die Messung der Inzidenzrate ist hilfreich, um die Dynamik der Prävalenzrate zu beschreiben, die sich aus der Bilanz der neuen Fälle und der Menge der die Prävalenz verlassenen Fälle (Heilungsrate, Letalitätsrate oder Fatalitätsrate), ergibt.

4.5.4 Die kumulativen Inzidenzraten, engl. cumulative incidence rate, CI

$$\text{Kumulative Inzidenzrate (CI)} = \frac{\text{neu Gefährdete, Exponierte oder Erkrankete } t_1 - t_0}{\text{nicht Gefährdete oder nicht Erkrankte } t_0}$$

CI schätzt die Wahrscheinlichkeit, dass sich eine gesunde Kohorte zum Zeitpunkt t_0 in einem bestimmten Zeitintervall nach der Einwirkung einer bestimmten Exposition infiziert bzw. erkrankt. Dieses Intervall entspricht also der Dauer, in der gesunde Menschen gefährdet sind, zu erkranken, bzw. dem Zeitintervall, in dem Menschen noch die Chance haben, gesund zu bleiben. Diese Wahrscheinlichkeit entspricht in ihrer formalen Konstruktion der, die in Sterbetafeln genutzt wird, um die verbleibende Lebenszeit nach der Geburt oder nach jedem weiteren Altersjahr zu messen. Dieser prinzipielle Aufbau kann auf jede Art von Ereignissen übertragen werden. Der begrenzende Faktor sind fehlende Daten. Aus diesem Grund wird CI meist unter realen oder quasi-experimentellen Studienbedingungen genutzt, die als reale oder fiktive Kohortenstudien konzipiert sind und die Schätzung von Ereigniswahrscheinlichkeiten ermöglichen.

Die Schätzung der Wahrscheinlichkeit für das Auftreten neuer Fälle zielt darauf ab, das Auftreten von Krankheiten in bestimmten Zeitintervallen nach einer Exposition einer gesunden Kohorte zu beurteilen. CI ermöglicht es also, die Anzahl neuer Fälle als Funktion der Zeit nach einer Exposition für vergleichbare Kohorten vorherzusagen. CI charakterisiert damit auch die Eigenschaften der epidemiologischen Potenziale.

Die kumulative Inzidenzrate ist also ein prädiktives Maß und zugleich ein Maß für die Bedeutung (Relevanz) einer Exposition in Bezug auf die weitere Lebenszeit nach der Exposition. Das Ergebnis ist offen für die Bewertung dieser Wahrscheinlichkeit entweder als Risiko oder als Chance mit einem Bereich zwischen 0 und 1. Die jeweils spezifischen CI unterschiedlicher epidemiologischer Potenziale können nicht addiert werden, etwa um die CI einer Gesamtbevölkerung in deren gesamter Lebenszeit zu ermitteln.

Die Messung von CI ist vor allem bei Kohortenstudien zu Krankheiten mit kurzen Intervallen zwischen Exposition und Auftreten von Effekten geeignet.

4.5.5 Inzidenzdichte (auch Force of Morbidity, momentane Inzidenzrate, Person-Zeit-Inzidenzrate, Hazardrate)

Die Inzidenzdichte misst die Auswirkungen einer Exposition als Funktion der Dauer unter Risiko.

$$\text{Inzidenzdichte (ID)} = \frac{\text{neue Fälle in } t + \Delta t}{\text{Personenzeit unter Risiko } t + \Delta t}$$

Bei dieser Kalkulation der Ereignisdichten wird die Anzahl der neu auftretenden Fälle als Funktion der Dauer einer Gefährdung geschätzt. Für diese Dauer gibt es zwei Möglichkeiten der Interpretation: Es kann sich einmal um die Zeitspanne zwischen einem einmaligen „Impakt“ (Verletzung, Applikation von Giften und Medikamenten etc.) handeln, oder um eine Zeitspanne kontinuierlicher oder diskontinuierlicher Schädigungen denen Menschen ausgesetzt sind. Im ersten Fall wird davon ausgegangen, dass die Anzahl der neu auftretenden Fälle nicht von der Wirkgröße der Schädigung ausgeht und sich im Ablauf einer realisiert, die dem Impakt eindeutig zugeordnet werden kann. Eine zweite Situation kann so beschrieben werden, dass eine Exposition dauerhaft kontinuierlich oder auch diskontinuierlich über längere Zeitspannen auf eine Person einwirkt. So entstehen für die Deutung der Messungen mindestens zwei unterschiedliche Konstellationen: (1) Eine Exposition ist schwerwiegend und führt nach vergleichsweise kurzer Zeit zu einer spezifischen Auswirkung (Erkranken oder Tod) und erreicht ihre maximale Ereignisdichte. (2) Eine Exposition ist niedrigschwellig und realisiert ihre Folgen erst nach langen Einwirkungsauern.

Das Konzept der Force of Morbidity wird verwendet, um die Intensität einer Exposition unter der Bedingung ihrer Zeitabhängigkeit zu messen. Dieses Maß ist die erste Ableitung nach der Zeit. Diese Rate wird auch als Hazard-Rate bezeichnet.

Die Inzidenzdichte gibt eine sinnvolle Information, wenn das Zeitintervall zwischen Einwirkung und Folge kurz ist. Es ist das Ziel, ein Maß für die "Kraft" oder für die Intensität einer definierten Exposition verfügbar zu haben, das die „Gefährlichkeit“ einer Exposition unabhängig von den Eigenschaften der von ihr betroffenen bewertbar

macht. Ein solches Maß soll also unabhängig von den Merkmalen der exponierten Personen sein und nur die Eigenschaften dieser Exposition beurteilbar machen. Es ist auch möglich, diese Dichte als Maß für die "Geschwindigkeit" zu interpretieren, mit der auf eine Exposition eine Wirkung folgt. Solche Intensitätsstudien setzen homogene Populationen voraus und sind eher für klinische als für bevölkerungsbezogene Studien geeignet.

Die Ereignisdichte ist kein Wahrscheinlichkeitsmaß für das Auftreten einer Krankheit, wie es die kumulative Inzidenzrate ist. Es ist auch zu beachten, dass wiederholbare Krankheiten als neue Fälle gezählt werden, bzw. eine Person mehrere Fälle verursachen kann. Das erschwert die Interpretation. Der reziproke Wert der Inzidenzdichte ist die durchschnittliche Risikozeit zwischen Exposition und Folge.

Die Dichtemaße unterstellen, dass ein expositionsspezifisches Folgeereignis ausschließlich von der Expositionszeit (time at risk) abhängt. Bei langen Zeitintervallen ist die im Grunde deterministische Konstruktion problematisch, da konkurrierende Risiken und die Modifikation der Betroffenen (z.B. Alterung) die Ergebnisse beeinflussen, ein Sachverhalt, der bei gutachtlichen Stellungnahmen erhebliche Bedeutung erlangen kann. Ein aus der Demografie stammendes Beispiel für ein Dichtemaß ist die sog „normale Lebensdauer“ (siehe Heft 2).

Anwendungsbeispiele aus der Industrie, dort z.B. zur Messung von Produktqualitäten benutzt, können dieses Dichte- oder Intensitätsmaß ebenfalls transparenter machen.

Beispiel:

Die von der Firma "A" hergestellten Glühbirnen können eine beworbene Lebensdauer von fünf Jahren haben, während das vergleichbare Produkt der Firma "B" eine Lebensdauer von drei Jahren hat. Die Verwendung dieser Daten für Qualitätsbeurteilungen ist sinnvoll, wenn man davon ausgeht, dass das bevorzugte Produkt tatsächlich die versprochene höhere Lebensdauer bei minimaler Varianz der Produkte hat. Eine weitere Bewertung könnte helfen, die Preise der Produkte an ihre unterschiedlichen Qualitäten anzupassen.

Ähnliche Beispiele finden sich in der Krankenversicherungsindustrie, wenn diese risikoäquivalente Produkte anbieten. Sie spielen auch bei Haftungskonflikten und gutachtlichen Stellungnahmen eine Rolle.

4.5.6 Heilungs- oder Genesungsraten (engl. recovery rate)

Kumulative Heilungsrate (engl. cumulative recovery rate, CRR)

$$\text{Kumulative Genesungsrate} = \frac{\text{genesene Personen } t_1 - t_0}{\text{kranke Personen in } t_0}$$

Diese CRR ist ein Wahrscheinlichkeitsmaß, das in seiner Konstruktion und seinen Interpretationsmöglichkeiten mit der kumulativen Inzidenzrate vergleichbar ist. Die Anpassung dieser CRR an die Sterbetafelmodelle ermöglicht auch die Modellierung der Wirkung von Epidemien durch die Konstruktion von Dekrementtafeln.

Rückfalldichte, auch Wiederholungsdichte

$$\text{Wiederholungsdichte} = \frac{\text{Genesene in } t+\Delta t}{\text{Personzeit unter Wiederholungsrisiko } t+\Delta t}$$

Dieses Maß ist formal der Inzidenzdichte vergleichbar. Sie wird verwendet, um die Geschwindigkeit zu beschreiben, mit der genesene Personen erneut erkranken oder einen Rückfall erleiden.

4.5.7 Quantifizierung von Sterbeereignissen

Die Mortalität dient der Beurteilung der krankheits- und populationsspezifischen Relevanz von Krankheiten.

In Staaten mit hinreichend korrekten Beurkundungen von Sterbefällen, von deren Ursachen und der erreichten Lebensalter erlauben solche Daten die quantitative Beurteilung der Gefährdungen durch Tod.

4.5.7.1 Rohe Sterblichkeitsrate

$$\text{rohe Sterblichkeitsrate} = \frac{\text{Anzahl der Gestorbenen } t_1-t_0}{\text{mittlere Anzahl des Potenzials } t_1-t_0}$$

Die Sterblichkeit bezeichnet die absolute Zahl der Personen, die im Beobachtungszeitraum sterben. Wenn sie für Gesamtbevölkerungen berechnet wird, ist die „mittlere Bevölkerung“ der Denominator. Diese Bezugsbevölkerung ist ein Querschnitt durch die in der Bevölkerung vertretenen Geburtsjahre zur Mitte des Jahres. In einigen Staaten ist sie der arithmetische Mittelwert der Summe der Bevölkerungsbestände zum Beginn und zum Ende des Beobachtungsjahres. Die Sterblichkeitsrate ist eine Anteilsziffer und kein Gefährdungsmaß. Für weiter Erläuterungen sei auf das Heft 2 dieses Kompendiums verwiesen.

4.5.7.2 Standardisierte Sterblichkeitsraten (SMR)

Die standardisierten Sterbeziffern zählen zu Vergleichszwecken die Anzahl der Todesfälle unter der Annahme einer standardisierten Bevölkerungsstruktur. Es handelt sich um eine strukturelle Bereinigung, die vorgenommen wird, um den Einfluss der Altersstruktur der Bevölkerung auf die Zahl der Todesfälle zu eliminieren. Obwohl es viele weitere Strukturmerkmale gibt, werden typischerweise nur die Altersstrukturen standardisiert. Deshalb sind die so genannten standardisierten Raten keine wirkliche Standardisierung von Auswirkungen der strukturellen Verschiedenheiten der zu vergleichenden Raten. Sie standardisieren allein die kalendarischen Altersstrukturen. Aus diesem Grund sind auch standardisierten Raten vorsichtig zu interpretieren, wenn sie als Argumente zur Erklärung von Ähnlichkeiten oder Ungleichheiten oder deren Veränderungen im Zeitverlauf herangezogen werden. Es muss auch berücksichtigt werden, dass das Alter kein stabiles biologisches Merkmal ist und innerhalb von Populationen und zwischen ihnen variiert und sich im Laufe der Zeit verändert.

4.5.7.3 Standardisierte Sterblichkeitsquoten (*standardized mortality ratio*)

Bei dieser Art von Messungen handelt es sich um eine Verhältniszahl von zwei standardisierten Sterblichkeitsraten. Die eine ist die beobachtete, aber standardisierte Rate und die andere eine erwartete standardisierte Rate. Die erwartete Rate ergibt sich aus der Annahme, dass die Zahlen aus einer stabilen Population als Maßstab für die folgenden Beobachtungszeiträume genommen werden können. Bei kurz andauernden Epidemien/Regressionen ist dieses Verhältnis hilfreich, um das Ausmaß der jeweiligen Wirkungen zu beurteilen.

4.5.7.4 Lebenszeitadjustierte Sterblichkeitsraten

Diese Raten berechnen die Sterblichkeit unter den Rahmenbedingungen einer Sterbetafel, bei der die Sterblichkeitsrate pro 1.000 der Tafelbevölkerung der reziproke Wert der Lebenserwartung ist. Dies ist eine Art standardisierte Sterblichkeitsrate, die an die Lebenserwartung einer Kohorte angepasst wird. Unter der Voraussetzung, dass die Lebenserwartung einer Geburtskohorte 100 Jahre beträgt, wäre diese Sterblichkeitsrate 10 pro 1.000. Sie wäre dann von anderen Einflüssen als der Sterbewahrscheinlichkeit nicht verzerrt. Wenn die Entscheidungsträger bei der Bekämpfung von Epidemien Prioritäten setzen wollen, könnten solche bereinigten Sterblichkeitsraten hilfreich sein.

4.5.7.5 Todesursachenziffern

Rohe ursachenspezifische Sterberate

$$= \frac{\text{Gestorbene an einer Todesursache in } t_1 - t_0}{\text{mittlere Bevölkerung } t_1 - t_0}$$

Für diese Raten sind die gleichen Argumente wie für die rohe Sterblichkeitsrate geltend zu machen. Es handelt sich um Anteilsziffern und nicht um Wahrscheinlichkeiten, bzw. Gefährdungsziffern. Die Summe aller todesursachenspezifischen Sterbeziffern ist immer die rohe Sterberate. Obwohl keine Gefährdungsmaße, spielen solche Todesursachenziffern in der öffentlichen und in der politischen Diskussion sowie in wissenschaftlichen Arbeiten eine in diesem Sinne gedeutete erhebliche Rolle. Dieses Problem wird verstärkt, wenn in Zeitreihenstudien nicht beachtet wird, dass alle sog. rohen Sterbeziffern in Phasen der Zunahme der mittleren Lebensdauer von Bevölkerungen einem parabolischen Trend folgen.

4.5.7.6 Standardisierte Todesursachenquote

Bei dieser Art von Anteilsziffern werden zwei standardisierte Todesursachenziffern, also altersstandardisierte Gliederungsziffern der Sterblichkeit zueinander in Beziehung gesetzt. Diese Kennzahl erfordert analoge Kommentare, wie sie bereits für die standardisierte Mortalitätsquote diskutiert wurden.

4.5.7.7 Übersterblichkeitsquote (engl. excess mortality rate, auch excess mortality ratio)

$$\text{Übersterblichkeitsquote} = \frac{\text{Gestorbene } (t_x - t_{x+})}{\text{Anzahl der erwarteten Sterbefälle } (t_1 - t_0)}$$

Die Übersterblichkeitsraten werden verwendet, um die Anzahl der Personen zu vergleichen, die in zwei unterschiedlich exponierten epidemiologischen Potenzialen versterben. Sie werden vor allem in Zeitreihenstudien verwendet.

4.5.7.8 Rohe Übersterblichkeitsquoten (excess mortality ratio)

$$\text{rohe Übersterblichkeitsquot} = \frac{\text{crude mortality rate measured } (t_x - t_{x+})}{\text{crude mortality rate measured expected } (t_1 - t_0)}$$

Diese Mortalitätsrate setzt zwei rohe Raten zueinander in Beziehung. Wenn die Sterblichkeitsrate, die als Nenner verwendet wird, höher ist als die Referenzrate, wird dies als Übersterblichkeitsquote gewertet. Diese Konstruktion geht von der Annahme aus, dass der Nenner als Standard oder als Norm verwendet werden kann. Diese Bedingung muss zunächst jedoch geprüft werden.

4.5.7.9 Standardisierte Übersterblichkeitsquote (engl. standardized excess mortality ratio)

$$\text{Standardisierte Übersterblichkeit} = \frac{\text{Gestorbene } (t_x - t_{x+})}{\text{standardisierte Anzahl der Gestorbenen } (t_1 - t_0)}$$

Diese Art der Messung der Übersterblichkeit wird zum Vergleich zweier verschieden stark exponierter, aber altersstrukturell vergleichbarer Bevölkerungen zur Beurteilung von Unterschieden der Sterblichkeit verwendet.

4.5.7.10 Standardisierte Übersterblichkeitsquote

$$\text{Standardisierte Übersterblichkeit} = \frac{\text{standardisierte Sterblichkeit der Zielbevölkerung } (t_x - t_{x+})}{\text{standardisierte Erwartungsrate } (t_1 - t_0)}$$

Der Unterschied zur Übersterblichkeit besteht in der Altersstandardisierung. Damit soll der mögliche systematische Fehler durch unterschiedliche Altersstrukturen vermieden werden.

4.5.7.11 Fatalitätsmessungen

Die Fatalität einer Krankheit bezeichnet die Wahrscheinlichkeit ihres tödlichen Ausgangs. Sie ist unter der spezifischen Bedingung zu erörtern, dass die Wahrscheinlichkeit zu sterben 1 ist. Die Lebenszeit ist also das wahre Maß für die Bewertung der Fatalität des Lebens und die Fatalität einer einzelnen Todesursache gleichsam eine Gliederung der Lebensrisiken. Die einzige Möglichkeit, diese "Risiken" zu bewerten, besteht darin, die unterschiedlichen altersspezifischen Verteilungen dieser Risiken innerhalb von Populationen zu vergleichen. Die Verteilung der Lebenszeiten ist veränderbar. Das ist nur durch die Veränderung der Fatalitäten der Todesursachen möglich. Die vielen möglichen Todesursachen "konkurrieren" um die Chancen, alt zu werden. Die Hoffnung ist immer, dass das Sterbealter in der Nähe des unbekannten biologischen Potenzials der Lebenszeit liegt. Wenn die mittlere Lebensdauer einer Bevölkerung niedrig ist, geht es vor allem um den Gewinn an Lebensjahren (adding years to life). In Bevölkerungen mit einer hohen mittleren Lebensdauer wird die Gewinnung von Lebensqualität zu einem wichtigen Ziel (adding life to years).

$$\text{Kumulative Fatalitätsrate} = \frac{\text{Gestorbene an einer Krankheit in } t_1 - t_0}{\text{Kranke an der betreffenden Krankheit in } t_0}$$

Die kumulative Fatalitätsrate ist eine Wahrscheinlichkeitsziffer. Sie ist also ein krankheitsspezifisches Gefährdungsmaß im Kontext der verfügbaren Möglichkeiten der

Medizin die Fatalitäten von Krankheiten zu beeinflussen. Diese Rate ist für die Bewertung von Epidemien und deren Verläufe besonders wertvoll. Sie ist zudem, leicht an die methodischen Bedingungen von Dekrementtafeln anzupassen und damit geeignet, die Auswirkungen einer Epidemie auf die mittlere Lebensdauer einer echten oder auch einer fiktiven Geborenenkohorte zu messen.

Die Summe aller kumulativen Fatalitätsraten (und spiegelbildlich der kumulativen Genesungsraten) ist für die Summe aller Todesursachen immer 1.

$$\text{Fatalitätsdichte} = \frac{\text{Gestorben an der Zielkrankheit } t + \Delta t}{\text{Personenrisikojahre für die Zielkrankheit } t + \Delta t}$$

Für diese Messung gelten die gleichen Anmerkungen wie für die oben beschriebene force of morbidity. Es ist notwendig, darauf hinzuweisen, dass diese Rate keine Schätzung der Überlebenszeiten durch die Kombination dieser Rate mit Dekrementtafeln ermöglicht.

$$\text{Letalitätsrate} = \frac{\text{Gestorbene an der Zielkrankheit } t_1 - t_0}{\text{mittlere Prävalenz der Zielkrankheit } t_1 - t_0}$$

Die Letalitätsraten sind „rohe“, bzw. nicht standardisierte Anteilsziffern. Im Falle von Infektionskrankheiten wird diese Rate sachlich falsch auch als "Infection Fatality Rate (IFR)" bezeichnet (*CFR, IFR, and You: Was ist die wahre COVID-19-Todesrate? | von The Blabbering Biologist | Medium; zuletzt geprüft am 01.08.2021*). Tatsächlich ist die sog. IFR eine Letalitätsrate. Das Problem ist, dass die Letalitätsraten einer Krankheit in der Regel niedriger sind als die Fatalitätsraten. Diese Verwechslung kann zu schwerwiegenden Fehlinterpretationen von empirischen Studien führen. Sie quantifiziert allein die relative Häufigkeit der Personen, die aufgrund ihres Todes aus der Prävalenz einer bestimmten Krankheit ausscheiden. Sie ist folglich kein adäquates Maß, wenn es darum geht, die durch die betreffende Krankheit bedingten Gefährdung durch Tod zu beurteilen.

Standardisierte Letalitätsrate

Dies ist die gleiche Rate wie die Letalitätsrate, allerdings an eine Standardbevölkerung angepasst, um Fehlinterpretationen aufgrund unterschiedlicher Altersstrukturen der jeweiligen Prävalenz der Erkrankten zu vermeiden. Die Verwendung dieser Art von Raten ist besonders problematisch, wenn die verglichenen Nenner klein sind und wenn die Todesursachen mit anderen relevanten Todesgefahren konkurrieren, wie es oft im hohen Alter der Fall ist.

Beispiel:

In der Bevölkerung X werden im Zeitraum t_{0-1} pro 100.000 Einwohner 20 neue Fälle von Lungenkarzinom festgestellt. Diese Aussage ist als logisch korrekt zu unterstellen, da die 20 neuen Fälle auch jeweils 20 verschiedenen Personen entsprechen werden.

Im Jahr 2002 sind in Deutschland 206.000 Frauen und 218.250 Männer an einer der Krebsarten (C 00- C 07 nach ICD) neu diagnostiziert worden. Bezogen jeweils auf alle Frauen, bzw. Männer entspricht dies einem Anteil von rund 0,49 und 0,54%, bzw. 488,5 bzw. 541,4 je 100.000 Frauen bzw. Männern (Krebsregister Saarland 2006). Diese Aussage ist nur dann korrekt, wenn keine der Personen zugleich an zwei oder mehr Karzinomen als neue Erkrankung diagnostiziert wurde.

Die altersspezifische Inzidenzrate gibt an, wie viele neue Fälle in einem bestimmten Altersabschnitt gezählt werden. Auf 100.000 Frauen der Altersgruppe werden bis zum 45. Lebensjahr 27, zwischen 45 und unter 60 Jahren 220, zwischen 60 und unter 75 Jahren 280 und jenseits dieses Alters 270 Frauen erstmals mit einem Brustkrebs diagnostiziert.

5 Relationale Häufigkeiten in der Epidemiologie

5.1 Das relative Risiko

Das relative Risiko ist der Quotient zweier Wahrscheinlichkeiten. In epidemiologischen Studien schätzt es das Ausmaß des Unterschieds zweier kumulierter Inzidenzraten in Bezug auf das die verglichenen Gruppen jeweils diskriminierende Merkmal. Es schätzt also Differenz von zwei disjunkten Ereigniswahrscheinlichkeiten.

Die Differenz zweier absoluter Risiken ist von der Auswahl der verglichenen Risikoklassen abhängig. Es kann mit relativen Risiken also nur eine relative Verschiedenheit beschrieben werden, die allein von der Auswahl des verglichenen epidemiologischen Potenzials abhängt. Über die Ursächlichkeit des Unterschiedes oder die Relevanz dieser Verschiedenheit in Bezug auf die Gesamtbevölkerung lassen sich keine Aussagen machen. Bei der Nutzung eines relativen Risikos krank zu werden, ist folglich die Risikobewertung ein Schlüsselproblem. Dies gilt sowohl in qualitativer wie in quantitativer Hinsicht, also sowohl bezüglich der Höhe des Unterschiedes als auch hinsichtlich seiner bevölkerungsbezogenen epidemiologischen Relevanz. So können geringe relative Risiken hoch relevant aber auch praktisch irrelevant sein. Das gilt analog auch für hohe relative Risiken.

Die Schwierigkeit der qualitativen Bewertung relativer Häufigkeitsunterschiede besteht darin, dass es kein Argument gibt, warum es zwischen Menschen nicht auch Gefährdungsunterschiede geben sollte. Diese "Ungleichheit" ist Teil der Individualität. Solange also Menschen Individuen sind, wird sich ihre Verschiedenheit auch messen lassen, und es werden nahezu beliebig viele relative Unterschiede beschreibbar sein.

Die Frage nach dem „richtigen“ Leben ist durch die Epidemiologie normativ nicht zu beantworten, nicht zu kontrollieren und auch nicht zu sanktionieren. Hierzu bedarf es externer Bewertungssysteme in Form formeller rechtlicher oder informeller sozialer Normen.

Das relative Risiko kann bei der Entscheidung hilfreich sein, ob ein hypothetischer Zusammenhang bestätigt werden kann oder nicht. Es ist jedoch selbst kein Gefährdungsmaß.

Trinkverhalten	Tod	Beobachtungsdauer in Personenjahren	Todeshäufigkeit pro 1000 Personenjahre	relatives Risiko
abstinenter	61	2083,2	25,3	1,00
Ex-Trinker	27	790,7	27,8	1,00

Typisch	89	5876,4	16,2	0,66
Alkoholi- ker	35	993,2	39,4	1,38

Tab. 7: Sterbefälle (alle Ursachen), differenziert nach dem Umgang mit Alkohol, modif. nach: Miller, G.J. et al, in: Inter. J. of Epidemiology, 19:1990, p. 928

Das Beispiel zeigt für die vier verglichenen Gruppen unterschiedliche relative Sterberisiken. Ob diese aber aus dem Alkoholkonsum ursächlich zu erklären sind oder ob sie nur allgemein auf eine strukturelle Verschiedenheit der Gruppen verweisen, kann aus den relativierenden Vergleichen allein nicht erklärt werden.

5.2 Odds und Odds Ratio

Odds bezeichnen Fallmengen, die zu einer abhängigen Variablen in Beziehung stehen.

Odds ratios entsprechen den Quotienten (Quoten) zweier Odds Ratios aus disjunkten Studiengruppen.

Die Odds, bzw. die Odds Ratios sind Assoziationsmaße, die vor allem in deskriptiven Querschnitts- und Prävalenzstudien genutzt werden. Ihr Wertebereich liegt zwischen 0 und unendlich. Im Gegensatz zum relativen Risiko beziehen sich die Odds Ratios nicht auf Ereigniswahrscheinlichkeiten.

Typischerweise werden Fälle oder Personen, die ein interessierendes Merkmal tragen, mittels sog. Vierfeldertafeln zueinander in Beziehung gesetzt, um die Merkmalsverteilungen zu gewichten, bzw. um die Ungleich- oder die Gleichverteilung der Merkmale über die alternativen Gruppen zu ermitteln. Das Ergebnis informiert über Assoziationsstärken der Merkmale, nicht über Kausalitäten oder bevölkerungsbezogene Relevanzen. Odds sind selbst keine Wahrscheinlichkeiten, informieren aber über die Chance unter einer Menge von „ähnlichen“ Personen/Fällen eine definierte Zielkrankheit häufiger als in den Vergleichsgruppen vorzufinden. Es handelt sich dabei praktisch immer um prävalente Fälle. Das Konzept ist für Suchstrategien, z.B. mit präventiven Zielsetzungen gut geeignet.

Odds Ratio	Relation der Fälle von Exponierten und Nicht-Exponierten ./. Relation der Nicht-Kranken unter den Exponierten und Nicht- Exponierten	$\frac{a/b}{c/d} = \frac{a \times d}{b \times c}$
---------------	--	---

(Synonyme in englischer Sprache: relative odds, cross-product ratio, exposure-odds ratio, disease-odds ratio, prevalence-odds ratio.)

Beispiel:

In einer Studie wird ermittelt, dass von 1500 untersuchten Kindern 150 behandlungsbedürftig krank sind. Bei der Aufgliederung dieser Kinder nach dem Sozialstatus der Eltern in zwei Gruppen werden 1300 als Kinder mit gutem bis sehr gutem Sozialstatus klassifiziert und 200 mit schlechtem Sozialstatus. In beiden Gruppen werden je 75 behandlungsbedürftige Kinder gefunden. Die Odds beträgt in den beiden Gruppen also einmal 75 zu 125, bzw. 75 zu 1225. Daraus folgt für die Chance ein behandlungsbedürftiges Kind in den beiden verglichenen Gruppen zu finden, ein Verhältnis von 10:1.

Die beschriebene Studiensituation ist die einer typischen Fall-Kontroll-Studie (engl. case-control-study): Es wird eine quantitative Relation, hier zwischen sozialer Situation und Behandlungsbedarf mit den Vorteilen/Nachteilen für jede der Kindergruppen beschrieben. Es gibt also einen nachvollziehbaren Anlass, eine Studie zu planen, um diese Beobachtung ursächlich aufzuklären, oder zumindest der benachteiligten Kindergruppe eine höhere Aufmerksamkeit zu widmen.

Beispiel:

Es wird in einer Untersuchung gefunden, dass von 100 Rauchern (Exponierte) 60 eine chronische Bronchitis haben und 40 nicht. Die Odds beträgt für die Raucher 60:40, bzw. 1,5. Analog wird für die Nichtraucher ein Verhältnis von 10 zu 90 ermittelt, also eine Odds von rund 0,11. Die Odds Ratio beträgt dann folglich 13,5. Neben offensichtlichen unmittelbar praktischen Konsequenzen wird offensichtlich eine Assoziation zwischen dem Rauchen und einer chronischen Bronchitis beschrieben. Es gibt darüber hinaus auch gewichtige Gründe, in weiteren Studien die Ursächlichkeit dieser Differenz zu untersuchen.

5.3 Das attributierbare Risiko

Das zuschreibbare oder attributierbare Risiko (attributable risk) informiert über den Anteil der Kranken gleicher Diagnose, deren Erkrankung einer spezifischen Variablen, u. U. also einer spezifischen Exposition ursächlich zuzurechnen ist.

Das attributierbare Risiko wird aus der Differenz der Inzidenz in der Gesamtbe-

völkerung und einer definierten Teilbevölkerung, bzw. einem definierten epidemiologischen Potenzial berechnet:

attributierbares Risiko	Anteil der Fälle einer Teilbevölkerung an den Fällen in der Gesamtbevölkerung	$I_g - I_e$
-------------------------	---	-------------

In Umkehrung wird auch versucht, das attributierbare Risiko als Maß für die zu erwartenden Präventionseffekte bei Ausschluss des Beitrags der Fälle der Teilbevölkerung zu deuten. Es spielt auch bei den Risikoselektionskonzepten von Äquivalenzversicherungen eine wichtige Rolle.

Das Konzept folgt der Annahme, dass sich die Gesamtmenge inzidenter Fälle aus der Addition zurechenbarer Teilrisiken ergibt.

	Lungenkrebs	koronare Herzkrankheit
Raucher	140	669
Nichtraucher	10	413
relatives Risiko	14,0	1,6
attributierbares Risiko	130	256

Tab. 8: Relatives und attributierbares Mortalitätsrisiko für Lungenkrebs und koronare Herzkrankheit bei Rauchern in einer Kohortenstudie unter britischen Ärzten, jährliche Mortalitätsrate pro 100.000 (zitiert nach: R. Doll and R. Peto, *Mortality in relation to*

6 Fehler in epidemiologischen Studien

6.1 Confounder

Confounder sind Wirkfaktoren, die selbst nicht Gegenstand einer Studie sind, auf die untersuchten Wirkungen aber Einfluss haben. Sind sie unbekannt oder nicht kontrolliert, können sie Studienergebnisse zur Wirkung der tatsächlichen Zielvariablen verfälschen und so zu Trugschlüssen führen.

Confounder sind reale Wirkgrößen, die jedoch die Analyse in Bezug auf eine andere Zielgröße beeinflussen. Sie sind also nur im Kontext einer spezifisch untersuchten Variablen auch „störende“ Variablen. Es sind besondere methodische Prozeduren erforderlich, um solche Confounder zu kontrollieren oder auszublenden. Es ist allerdings in der Regel unerlässlich, sie bei der Bewertung eines Forschungsergebnisses für die Handlungspraxis dann wieder in Betracht zu ziehen.

In diesem Sinne kann in dreierlei Beziehung von Confounder gesprochen werden:

1. Zwei (bzw. mehrere) Ursachen erzeugen gemeinsam eine Wirkung, die dann irrtümlich jedoch nur einer Ursache zugeschrieben wird.
2. Eine Wirkung ist das Ergebnis keiner der untersuchten Ursachen, wird denen aber irrtümlich zugeschrieben.
3. Von Effekt-Modifikation wird hingegen gesprochen, wenn eine Wirkung richtig einer Ursache zugeschrieben wird, aber durch andere ursächliche Einflüsse in der Wirkrichtung oder -intensität verändert und bei Nicht-Beachtung dieses Zusammenhanges fehlerhaft bewertet wird.

Die Problematik der Confounder wird dadurch verstärkt, dass zu Beginn einer Studienserie oft nicht klar ist, welche der außerhalb des Studienziels liegenden Einflüsse das Ergebnis verzerren können.

Eine gründliche Diskussion möglicher confounder ist methodische Norm jeder epidemiologischen Studie.

6.2 Bias

Bias sind Ursachen einer Abweichung des Forschungsergebnisses von der Realität.

Bias entstehen auf vielfältige Weise: bei der Datengewinnung, durch die Veränderung der untersuchten Grundgesamtheit durch das Studiendesign, durch die

Datenauswertung, durch die Interpretation, die Art der Veröffentlichung, durch die Zitierweise der Studienergebnisse usw. Es handelt sich bei ihnen um systematische und unsystematische Fehler. Vor allem letztere lassen sich praktisch nie völlig kontrollieren. Fehlt in einer epidemiologischen Studie die vollständige Dokumentation des Ablaufs einer Studie und ihrer Datenquellen, lassen sich Bias in der Regel nicht mehr rekonstruieren.

Da bestimmte Bias immer wiederkehren bzw. regelhaft zu beachten sind, hat sich für sie eine eigene Begrifflichkeit herausgebildet.

Typische und häufige Bias sind:

1. Informationsfehler (engl. information bias),
2. Selektionsfehler (engl. selection bias),
3. Ergebnisverzerrungen durch die Unterschiedlichkeit der Zeitpunkte zu denen eine Krankheit diagnostiziert wird (engl. lead-time bias),
4. Veränderung des Studiensamples durch unkontrollierte Verluste von Probanden (engl. attrition bias)
5. Fehler durch Erinnerungslücken von anamnestic Befragten (engl. recall bias),
6. Bewertungsfehler (engl. detection bias)
7. Klassifikationsfehler (engl. misclassification bias).
8. Weitere zu beachtende Bias sind z.B.: response bias, sampling bias, ascertainment bias, observer bias, interviewer bias, measurement bias, data presentation bias, in assumption bias, interpretation bias, publication bias, reader bias.

7 Der Zusammenhang zwischen Inzidenz und Prävalenz

Die Prävalenz erneuert sich in der durchschnittlichen Dauer der Fälle zwischen Zugang und Ausscheiden aus der Prävalenz.

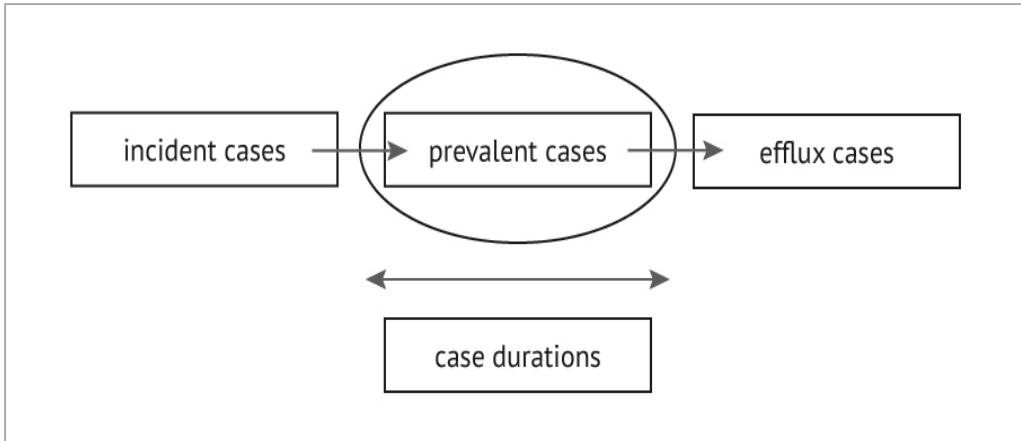


Abb. 7: Die Erneuerung der Prävalenz ((Niehoff, J.U., Li, B. *Epidemics and Regressions*. Saluscon Akademie Verlag, Berlin 2022, ISBN 978-3-948267-15-5)

Die Menge der neuen Fälle (Inzidenz) entspricht der Wahrscheinlichkeit, richtig positiv diagnostiziert zu werden multipliziert mit der Menge, die Zugang zur Diagnostik hat oder die Möglichkeit einer Diagnostik aktiv nutzt oder von verpflichtenden Testungen erfasst wird. Im Falle übertragbarer Krankheiten sind zudem die Dynamik der Potenziale der Infizierten und die der Kranken und die jeweilig zugehörigen Bewegungsmengen strikt zu unterscheiden.

Epidemien, die ursächlich auf die Zunahme von Gefährdungen zurückgehen, folgen typischen zyklischen Verläufen, die nach Niehoff und Li in 8 Phasen unterschieden werden können (siehe (Niehoff, J.U., Li, B. *Epidemics and Regressions*. Saluscon Akademie Verlag, Berlin 2022, ISBN 978-3-948267-15-5).

Die Menge der prävalenten Fälle ist dem Produkt der Diagnosewahrscheinlichkeit und der Behandlungsdauer proportional. Ob und ggf. in welchem Ausmaß die Mengen diagnostizierter Fälle mit denen der neuen Erkrankungen übereinstimmen, bleibt immer im Einzelfall der Bewegungen von Prävalenzen zu entscheiden.

Die Behandlungsdauer entspricht der Verteilung der Behandlungsdauern über die Menge derer, die Zugang zu einer Therapie haben. Sie bildet sowohl die patientenbezogenen Möglichkeiten, Fähigkeiten und Bereitschaften an der Therapie mitzuwirken ab, wie auch die medizinischen Interventionen und deren Zugänglichkeit ab.

Je nach Sachzusammenhang sind die Zugänge zur Prävalenz z.B. neuerkrankte oder neu diagnostizierte Fälle. Die Bewegung der Inzidenzen ist folglich entweder von den Erkrankungs-, von den Diagnosewahrscheinlichkeiten oder von beiden bestimmt. Analog können aber auch gesunde, bzw. an einer Zielkrankheit nicht kranke Personen als Zugänge zur Prävalenz der Gefährdeten bezeichnet werden und dann z.B. als das Potenzial unter Risiko klassifiziert werden. Entsprechend kann das Ausscheiden aus der Prävalenz Kranker Folge der Beendigung der Therapie wegen Genesung, wegen des Übergangs in den Zustand Behinderung oder auch wegen des Todes einer Person sein. Analog entspricht das Ausscheiden aus der Prävalenz der Gefährdeten dann der Menge derer, die Erkranken, bzw. der Menge derer, für die sich die Gefährdungslage verändert hat.

Wenngleich diese Erneuerungsprozesse der Prävalenz aus Mangel an entsprechenden Informationen nur selten dargestellt werden können, ist dennoch richtig, dass die Dauer, also z.B. die Behandlungsdauer, ein besonders wichtiges Maß für die Beurteilung der Wirkungen der medizinischen Versorgung auf die Prävalenz ist.

Diese Dauer ist vielfältig bedingt. Sie hängt einerseits von Patientenmerkmalen, wie dem Alter bei Behandlungsbeginn oder der Mitwirkungsbereitschaft und -fähigkeit der Patienten und deren Eigenschaften ab. Sie hängt aber auch von der Art und den Eigenschaften der Krankheitskrankheit selbst ab. Sie ist schließlich von allen Dimensionen der Qualität medizinischer Versorgung mitbestimmt.

Wenn angenommen wird, dass die Menge der neuen Erkrankungen und die Krankheitsdauer konstant sind, gilt für den Vorgang der Erneuerung der Prävalenz (Erneuerungsprozess) unter Gleichgewichtsbedingungen:

$$\text{Erkrankungswahrscheinlichkeit} \times \text{Potenzial der Gefährdeten} \times \text{Krankheitsdauer} = \text{Prävalenz}$$

Unter realen Bedingungen ist ein solches Gleichgewicht nur selten und nur zufällig oder nur für eine begrenzte Anzahl von Beobachtungsjahren zu erwarten. Ziele der Medizin sind schließlich, sowohl auf die Inzidenz und auf die Behandlungsdauer einzuwirken. Das Ungleichgewicht zwischen Zugängen und Abgängen ist also ein explizites Präventions- und Versorgungsziel und für viele Krankheiten die Regel. Da in der Regel die Behandlungsdauer eher ermittelt werden kann als die Krankheitsdauer, muss diese über die Behandlungsdauer geschätzt werden.

8 Der Vergleich in der epidemiologischen Forschung

Die Methodik(en) des Vergleichs von Daten sind die wichtigsten Erkenntnisquellen der Epidemiologie.

Die Arbeitsprinzipien vergleichender Studien in der Epidemiologie sind

1. die Formulierung der Fragestellung bzw. des Forschungsziels
2. die Formulierung von Hypothesen zu den Beziehungen zwischen unabhängigen und abhängigen Variablen sowie die Entwicklung von Hypothesen über Ursache-Wirkungs-Beziehungen
3. die formal korrekte Beschreibung und Erhebung von Gesundheitsdaten in Bezug auf die Fragestellung
4. die Beurteilung der Datenqualität
5. die Analyse von Fehlerquellen (bias, confounding)
6. die Bewertung von Ergebnissen in Bezug auf die Fragestellung und ggf. hinsichtlich ihrer praktischen Nutzung in Programmen des Gesundheitsschutzes und der medizinischen Versorgung.

Vergleiche in der Epidemiologie können die Erzeugung der Vergleichbarkeit (z.B. durch Datenstandardisierung und Matching) von Daten ebenso erfordern wie der explizite Nachweis der Unvergleichbarkeit der Daten ein primäres Forschungsziel sein kann.

Im „Ungleichheitsfall“ können die Vergleichsmethoden grundsätzlich nur zwei Formen annehmen: (1.) Gleiche Bevölkerungen werden bezüglich der Ursächlichkeit ungleicher Häufigkeiten einer Zielvariablen untersucht. (2.) Ungleiche Bevölkerungen werden bezüglich der Ursächlichkeit gleicher Häufigkeiten einer Zielvariablen untersucht.

Für diesen Sachverhalt unterschied der englische Ökonom John Stuart Mill (1806-1873) zwei Studienkonzepte, die er die Methode der Konkordanz oder Übereinstimmung und die Methode der Differenz nannte (vergl. Niehoff, J.-U. *Risikofaktoren, Risikofaktorentheorie, Risikokonzept* Z.ärztl.Fortb. Teile 1 und 2. 1978, 72:84-89; 72:145-149; Nohlen, D. (2004): *Vergleichende Methode*. in: Nohlen/Schultze (Hrsg.): *Lexikon der Politikwissenschaft*. 2.Aufl., Verlag Beck, München, S.1042-1052).

Für die Methode der Konkordanz müssen die untersuchten Bevölkerungen möglichst ähnlich sein. Bezogen auf medizinische Fragestellungen hieße dies, möglichst homogene Teilmengen zu erzeugen, um die Wirkung einer oder mehrerer Einflussvariablen prüfen zu können. Dieses Vorgehen entspricht z.B. den Prinzipien klinischer Studien.

Bei der Differenzmethode sind die untersuchten Einflussvariablen zwar regelhaft vielfältig, die Kontextfaktoren jedoch ähnlich. Es wird also einmal geprüft, warum ungleiche Einflüsse gleiche Wirkungen erzeugen und zum anderen, warum gleiche Einflüsse unterschiedliche Wirkung entfalten. Die erste der Methoden zielt auf die Messung spezifischer Wirkungen bei einer Einflussnahme. Die zweite Methode prüft gleichsam die Relevanz eines solchen Einflusses auf, im Falle der Epidemiologie, die gesamte Studienbevölkerung.

Beispiel:

In einer klinischen Studie (Methode der Konkordanz nach J.S. Mill) wird gezeigt, dass hoher Blutdruck das Risiko erhöht, einen Herzinfarkt zu erleiden. Hoher Blutdruck wäre also ein spezifischer Einfluss. Eine solche Studie setzt voraus, dass die Studiengruppe möglichst homogen ist und dann die Wirkungen (Infarkt) nur in Abhängigkeit vom Blutdruck variieren. In einer epidemiologischen Studie wäre hingegen mittels der Differenzmethode zu prüfen, ob, bzw. ggf. in welcher Bevölkerung dieses klinische Ergebnis ein relevanter Ansatz ist, um die Infarkthäufigkeit durch die Veränderung der Blutdruckwerte zu verringern. Es wäre also z.B. die Häufigkeitsverteilung der Blutdruckwerte in einer Bevölkerung zu ermitteln. Der Ansatz von Mill erklärt auch, warum die Wirksamkeit im Einzelfall und die Relevanz auf Bevölkerungsebene zwei unterschiedliche Fragestellungen sind und die Wirksamkeit im Einzelfall und die Relevanz für eine Gesamtheit von Elementen häufig nicht gemeinsam zu erreichen sind (J. St. Mill (1882) System of Logic, Vol. I, Buch 3, Kapitel 8: „Of the Four Methods of Experimental Inquiry“. Eighth Edition, New York, Harper & Brothers, Publishers, Franklin Square). Dieses Problem hat auch den britischen Internisten G. Rose (1926-1993) beschäftigt und ihn zu der konzeptionellen Position geführt, dass auf Bevölkerungsebene Präventionsprogramme zur Senkung der Infarkthäufigkeiten durch Programme der Blutdrucksenkung nur messbare Ergebnisse erzeugen können, wenn die Zielwerte für den Blutdruck in einem Maße abgesenkt werden, dass es auch messbare Präventionsergebnisse geben kann. Dies wäre ggf. auch dann zu fordern, wenn der weitaus überwiegende Teil der Bevölkerung falsch positiv behandelt würde, also auch dann, wenn die Wahrscheinlichkeit des Eintritts eines Infarktes gering ist. Diese Strategie nannte er „mass strategy“ (auch population strategy). Sie war für ihn die notwendige Folgerung aus dem „Versagen“ sog. „high risk strategies“ auf Bevölkerungsebene.

Der Vergleich epidemiologischer Daten, sei es in räumlicher, sozialer oder zeitlicher Hinsicht oder zwischen exponierten oder nicht-exponierten Personen usw. verlangt zunächst, die Vergleichbarkeit zu prüfen und ggf. herzustellen. Dabei geht es sowohl um die Vergleichbarkeit der gezählten Fälle (Zähler) im Sinne der richtigen diagnostischen Zuordnung als auch um die Vergleichbarkeit der Bezugsgesamtheit („Risikozeit“ der Gesamtbevölkerung).

Die Sicherung der Vergleichbarkeit des "Nenners" ist in der Regel schwieriger als die des "Zählers". Diesem Sachverhalt muss in der Forschungspraxis eine hohe Aufmerksamkeit geschenkt werden.

In epidemiologischen Bevölkerungsstudien ist die Heterogenität der Bezugsbevölkerung ein grundlegender Sachverhalt. Anders als in klinischen Fallstudien, in denen z.B. Aussagen zu therapeutischen Wirkungen angestrebt werden und deshalb die Homogenität der Studienpopulation gesichert werden muss, zielen epidemiologische Studien auf Kenntnisse über Bevölkerungen, also z.B. auf die Offenlegung ihrer Heterogenität. Diese ist dann auch keine „Störgröße“, sondern das Analyseziel.

Vergleiche spielen in der Epidemiologie eine weitaus größere Rolle als z. B. experimentelle Studien. Sie sind zudem häufig auch zu problematisieren, wenn sie, wie im Falle von Interventionsstudien Bevölkerungen zu ihren Objekten haben.

8.1 Die Vergleichbarkeit des gemessenen Nominators

Für die Zuordnung eines Untersuchungsgegenstandes sind verbindliche, reproduzierbare und logisch sinnvolle Zuordnungsregeln erforderlich, z.B. in Form gleichartiger Definitionen von Krankheit. Dabei ist ggf. die „Gleichartigkeit“ wichtiger als die „Richtigkeit“.

Hierzu dienen u.a.

- standardisierte Untersuchungstechniken und Diagnoseverfahren
- einheitliche Krankheitsbegriffe oder Normwerte
- verbindliche und logisch nachvollziehbare Standards bei der Bildung von Messwertkategorien

8.2 Die Vergleichbarkeit des Denominators

In der Epidemiologie sind die Denominatoren die epidemiologischen Potenziale.

Für epidemiologische Studien ist zu fordern, dass die Referenzbevölkerung in der Varianz ihrer epidemiologischen Potenziale darstellbar wird. Jedes der Potenziale fasst ähnliche, also nicht gleiche oder identische Elemente zusammen. Solche Ähnlichkeiten zielen primär auf die Sinnhaftigkeit von Handlungskonzepten für Menschen, deren Lebenskontexte hinreichend ähnlich sind (Settings). Unübersehbar ist allerdings ein Problem, dass die Bestimmung des logisch zutreffenden Denominators die Kenntnis voraussetzt, dass die Nominatoren zu ihm tatsächlich in einer ursächlichen Beziehung stehen. Eine solche Kenntnis kann nicht immer vorausgesetzt werden. Folgerichtig folgen epidemiologische Studien in einer großen Zahl Suchstrategien, die nicht

zu einer Entscheidung über die Annahme oder die Ablehnung einer Hypothese führen, sondern der Begründung von Zusammenhangshypothesen dienen, die dann mit weiteren Methoden und Studienkonzepten zu prüfen sind.

Für die Sicherung der Vergleichbarkeit der epidemiologischen Potenziale gibt es allgemein drei Verfahren:

1. die Selektion, um entweder gleichartige oder gleichstrukturierte Bezugsdaten nutzen zu können (Matching)
2. die Standardisierung und Klassifizierung
3. die Randomisierung

Zwischen beiden Verfahren liegt ein weites Spannungsfeld der Interpretation. Die Homogenisierung des Bezugssystems schafft die Vergleichbarkeit, bzw. strukturelle Ähnlichkeit der Bezugsmengen, um Hypothesen hinsichtlich der Wirkung bestimmter vermuteter Faktoren prüfen zu können. Für die Beurteilung der Relevanz solcher Faktoren für das praktische Handeln ist dann jedoch immer die reale Inhomogenität der Bezugsbevölkerung zu beachten.

9 Die epidemiologischen Bewegungsformen

Die Bewegungsformen der Prävalenz folgen aus ihren Ursachen.

Die Bewegungsformen der Prävalenz werden nachfolgend in einer schematischen Übersicht zusammengefasst. Zum Zwecke der Vereinfachung wird diese für die Prävalenz Kranker dargestellt. Das Modell lässt sich aber auch an die Prävalenzen jedes anderen gesundheitsrelevanten Zustands anpassen.

Bezeichnung der Veränderung	Unabhängige Variable		Abhängige Variable	
	Inzidenzfälle	Dauer	Prävalenz	Abgänge
(1) Gleichgewicht	==	==	==	==
(2) Epidemie		==		
(3) Regression		==		
(4) Therapeutisch bedingte Epidemie	==			
(5) Therapeutisch bedingte Regression	==			
(6) Diagnostisch bedingte Epidemie				==
(7) Diagnostisch bedingte Regression				==

Zeichenerklärung:

==

keine Veränderung

nachfolgendes Ansteigen

nachfolgendes Absinken

Ansteigen

vorübergehendes Ansteigen

Absinken

vorübergehendes Absinken

Abb. 8: Die Bewegungsformen der Prävalenz und ihre Ursachen (nach B. Kreuz und Mitarbeiter: Über den Zusammenhang zwischen Krankheitsdauer und Morbidität. Ztschr. ärztl. Ausbildung 67 (1973) 144-148)

Die Menge der zu einem Zeitpunkt vorhandenen epidemiologischen Fälle hängt ab von

- der Menge der inzidenten Fälle
- der Menge der Fälle, die aus der Fallmenge ausscheiden
- der durchschnittlichen Dauer zwischen der Inzidenz und den Fällen, die aus dem Bestand der Fälle ausscheiden, also, sofern eine Erkrankung auch zu einer Behandlung führt, von der Behandlungsdauer.

Die Menge der Zugänge, bzw. die Inzidenz kann verändert werden

- durch Veränderungen der force of morbidity
- durch Veränderungen der Prävalenz der Exponierten
- durch Veränderungen der Zeitpunkte der Diagnosestellung (Verkürzung/Verlängerung des diagnostischen Intervalls)
- durch Veränderung der Menge der als krank richtig positiv und falsch positiv diagnostizierten Fälle (z.B. durch den Wandel der diagnostischen Zuordnungskriterien, der sozialen Zugänglichkeit medizinischer Versorgung, durch den Wandel der Inanspruchnahme von Gesundheitseinrichtungen durch die Bevölkerung u.ä.m.)
- durch die Verkürzung oder Verlängerung der Behandlungsdauer (Wandel therapeutischer Möglichkeiten und Strategien, Wandel der Struktur der Behandelten)

Zur Vereinfachung seien diese Möglichkeiten hier als

1. Gefährdungsänderung
2. Diagnostikänderung
3. Therapieänderung

zusammengefasst.

Damit können alle möglichen Prävalenzänderungen (also die Bewegungsformen der Prävalenz) auf Veränderungen der Inzidenzfälle und/oder der Behandlungsdauer zurückgeführt werden.

Es lassen sich insgesamt sieben Bewegungsformen der Prävalenz unterscheiden. Diese sind

1. die risikobedingte Epidemie
2. die diagnostisch bedingte Epidemie
3. die therapeutisch bedingte Epidemie
4. die risikobedingte Regression
5. die diagnostisch bedingte Regression
6. die therapeutisch bedingte Regression
7. das Gleichgewicht

Diese Zusammenhänge sind erstmals von Bernhard Kreuz (1919 – 1980) im Ergebnis einer historischen Analyse der Prävalenzbewegung des Diabetes mellitus beschrieben und von seinen Mitarbeitern dann weiterentwickelt worden.

Die Bewegungen der Prävalenz sind keine Eigenschaften von Krankheiten, sondern von sozialen Prozessen. Eine reale Bevölkerung ist immer heterogen und die in ihr vorkommenden Krankheiten sind stets ungleich verteilt. Es können deshalb bei gleichen Krankheiten gleichzeitig unterschiedliche Bewegungen in Teilen der Bevölkerung nebeneinander ablaufen. Die Dynamik einer krankheitsspezifischen Prävalenz in der Gesamtbevölkerung ist folglich die Resultante von in Richtung und Intensität unterschiedlichen Bewegungen in den Teilbevölkerungen.

Aus analytischer Sicht ist es nicht allein bedeutsam die jeweils resultierende Bewegungsform der Fälle zu identifizieren, sondern vor allem jene Teilmengen einer Bevölkerung, die von ihr jeweils betroffen sind.

Die häufig zu findende Formulierung, eine Krankheit nähme zu oder auch ab, ist unsinnig. Selbstverständlich nimmt nicht die Krankheit zu oder ab, sondern die Anzahl der Fälle, ggf. auch die Menge der Menschen, die an ihnen leiden. Diese an sich triviale Richtigstellung ist dennoch notwendig. Sie verweist darauf, dass Aussagen der Epidemiologie nur dann sinnvoll sind, wenn auch bezeichnet werden kann, über wen oder was Aussagen gemacht werden. Dies sind Aussagen über Menschen, bzw. Gruppen von Menschen, keine Aussagen über Krankheiten. Eben hierin liegt aber auch eine der spezifischen Erkenntnisfunktionen der Epidemiologie.

9.1 Systematik der Ursachen von Bewegungen der Prävalenz

Alle Veränderungen der Prävalenz lassen sich vollständig auf vier Ursachen zurückführen.

Die Bewegungsursachen der Prävalenz sind

1. demographische Faktoren
2. medizinische Faktoren
 - a. diagnostische Faktoren
 - b. therapeutische Faktoren
3. epidemische Faktoren
4. präventive Faktoren

Allen dieser Faktoren ist gemeinsam, dass sie direkt und indirekt auf von Menschen verursachte Einflüsse zurückgehen, also soziale Ursachen sind.

Die Bewegungsursachen der Prävalenz wirken nie auf alle Personen und auf alle vorkommenden Krankheiten mit gleicher Intensität und Richtung. Deshalb können Teilgruppen einer Bevölkerung mit ansonsten gleichen Gesundheitsproblemen Prävalenzbewegungen aufweisen, die in Richtung und Intensität deutlich verschieden sind. Die Dynamik der Prävalenz einer Gesamtbevölkerung ist also immer der Mittelwert der Bewegungen in den relevanten Teilbevölkerungen, bzw. Potenzialen.

Aussagen zur Veränderung der Prävalenz einer Krankheit sollten stets mit der konkreten Bezeichnung der Bevölkerungsgruppen oder sozialen Schichten verbunden werden, die diese betreffen. Nur so werden sie auch handlungsrelevant. Unterschiedliche Prävalenzbewegungen in Teilbevölkerungen erzeugen auch die Unterschiede zwischen ihnen.

9.1.1 Demografische Faktoren

Demografische Faktoren verändern den Umfang und die Struktur einer Bevölkerung.

Die prinzipiellen Wirkungen demografischer Faktoren sind im Heft 2 dargestellt. Es kann davon ausgegangen werden, dass die wichtigsten Ursachen für wachsende medizinische Versorgungsbedarfe in Bevölkerungen bis in die jüngere Vergangenheit aus der Zunahme der mittleren Lebensdauer und aus Effekten des Strukturwandels von Bevölkerungen folgen. Im globalen Vergleich gilt dies mit großer Sicherheit auch weiterhin

Der quantitative Bedarfswandel geht also nicht primär auf Gefährdungszunahmen zurück, sondern auf die Zunahme der Menge von Personen, die den Gefährdungen ausgesetzt sind, sowie auf den Wandel der Bevölkerungsstrukturen.

9.1.2 Medizinische Faktoren

a) Diagnostik

Diagnostische Faktoren verändern die Menge der als Inzidenz klassifizierten Fälle.

Die Veränderung der Inzidenz durch diagnostische Effekte in aufeinander folgenden Zeitintervallen kann folgende Ursachen haben:

- Veränderung der Zugänglichkeit diagnostischer Möglichkeiten (mehr oder weniger Kranke haben die Chance, als krank erkannt zu werden)
- Veränderung von Normwerten
- Differenzierung von bislang als gleich gedeutete Krankheiten Zusammenführung von bislang als verschieden gedeutete Krankheiten
- Veränderung des Zeitpunkts der Diagnosestellung (Kranke werden früher oder später als krank erkannt)
- Ausweitung der „Nachfrage“ diagnostischer Aktivitäten z.B. als Folge von Marketingkonzepten (engl. supplier induced demand)

Entwicklungen in Richtung einer Individualisierung von Therapiekonzepten werden also immer auch, ggf. erhebliche, Folgen für die Prävalenzen von Krankheiten haben.

Eine Mengenausweitung der Diagnostik liegt vor, wenn der Zugang zur Diagnostik ausgeweitet wird und rechtliche oder ökonomische Regulationsmechanismen in der medizinischen Versorgungspraxis eine vermehrte Anwendung diagnostischer Verfahren induzieren (Überversorgung). Analog werden Mechanismen, die Unterversorgung erzeugen, zu einer Mengenreduktion führen.

Inzidenzänderungen infolge der zeitlichen Vorverlagerung von Diagnosezeitpunkten sind vorübergehender Art. Inzidenzänderungen infolge einer Ausweitung der diagnostischen Aktivitäten verbunden mit der Absenkung von Normen für die Definition von Krankheit und den Behandlungsbedarf streben einem dauerhaft höheren Plateau zu.

Veränderungen des Zeitpunktes der Diagnosestellung folgen

- neuen diagnostischen Techniken und Prinzipien und
- dem Wandel der therapeutischen Konzepte, ob und ggf. ab wann ein Zustand als krank, bzw. als behandlungsbedürftig/-fähig gilt, dieser Wandel schließt regelhaft eine veränderte Bereitschaft ein, falsch positive und falsch negative Diagnostikergebnisse zu akzeptieren.

In der medizinischen Versorgungspraxis überlagern sich die Mengenausweitungen und die Zeitverschiebung. Sie resultieren sowohl aus Innovationen wie aus wirtschaftlichen Interessen und werden zumeist als Präventionsangebot begründet.

Es wird bei solchen Entwicklungen zumeist in Kauf genommen, dass „Mengenauswei-

tungen“ ggf. erhebliche Rückwirkungen auf die Sensitivität und die Spezifität diagnostischer Tests haben und dann die Risiko-Nutzen-Relation ungünstig beeinflussen.

Vor allem eine Veränderung der Definitionskriterien (Normwerte) führt zu Zeitverschiebungen, die regelmäßig auch zu einer Strukturänderung der diagnostizierten Fälle etwa hinsichtlich der Merkmale Alter und Schweregrad führen. Sie können so auch die Risiken therapeutischen Intervention positiv wie auch negativ beeinflussen. Diese Entwicklungen vollziehen sich weltweit und haben besonders in Bevölkerungen, deren Gesundheitssicherungssysteme noch in frühen Entwicklungsphasen sind, erhebliche und zumeist massiv negative Folgen.

Änderungen der Diagnostikkriterien wirken immer auf den Umfang und die Struktur der Prävalenz zurück.

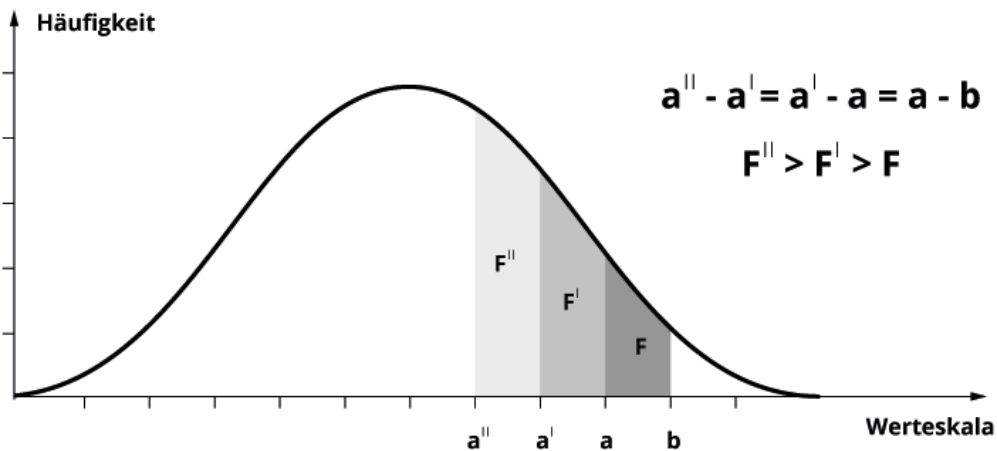


Abb. 9: Der Zusammenhang von Normwertänderungen und diagnostizierten Fällen

Krankheitsstadien 1 - 5 (1 = sehr früh; 5 = sehr spät)	Ausgangssituation hinsichtlich der Diagnostik von Krankheit bei t ₀	Effekt der Zeitverschiebung bei t ₁	Effekt der Zeitverschiebung bei t ₂
1	5	10	15
2	10	15	20
3	15	20	25
4	20	25	25
5	25	20	15
nie diagnostiziert	25	10	5
Zunahme der	um 0	um 20 % auf 90 % aller	um 27 % auf 95 % aller

Diagnostizierten in %		Kranken	Kranken
-----------------------	--	---------	---------

Tab. 9: Simulation der Wirkungen von Zeitverschiebungen auf die Struktur der Fälle nach dem Merkmal "Krankheitsstadium"

Bei der Interpretation des fiktiven Beispiels in der Tab. 5 ist zu beachten:

- Diagnostiziertenraten von 75, 90 oder 95 % sind nur bei wenigen und dann schweren Krankheiten anzunehmen
- eine Zunahme der Prävalenz um 20 % ist bereits erheblich
- bedeutsam ist vor allem der Strukturwandel der Diagnostizierten
- tendenziell werden die Effekte solcher Zeitverschiebungen kleiner bzw. gleiche quantitative Effekte lassen sich nur mittels wachsender Aufwände erzielen

b) Therapie

Therapeutische Faktoren verändern die Zeitpunkte und/oder die Art des Ausscheidens aus der Menge der behandelten Fälle. Änderungen der Therapie wirken immer auf den Umfang und die Struktur der Prävalenz zurück.

Neue oder veränderte Therapiestrategien können nur drei Wirkungen haben:

1. sie verändern die Lebensqualität
2. sie verändern die Verteilung der Behandlungsdauer unter den Therapierten
3. sie verändern die Arten des Ausscheidens aus der Prävalenz der Therapierten

Hinsichtlich der Arten des Ausscheidens sind die Alternativen „Heilung/Defektheilung“ oder Tod zu differenzieren.

Therapeutische Interventionen bei heilbaren Krankheiten sind darauf gerichtet, Krankheitsfolgen zu vermeiden und die Heilung möglichst schnell herbeizuführen.

Therapeutische Interventionen bei nicht heilbaren, bzw. noch nicht heilbaren Krankheiten sind darauf gerichtet, die Lebensdauer (damit zugleich meist auch die Behandlungsdauer) möglichst zu verlängern und/oder die Lebensqualität zu bessern.

Verlängerungen der Behandlungsdauer verzögern vorübergehend das Ausscheiden aus der Menge der Behandelten, Verkürzungen der Behandlungsdauer führen hingegen zu einer vorübergehenden Zunahme der aus der Menge der Behandelten ausscheidenden Personen.

9.1.3 Epidemische Faktoren

Die Ursachen von gefährdungsbedingten Epidemien bewirken die räumliche, und/oder zeitliche Ausbreitung von Gesundheitsgefährdungen (Expositionen).

Eine durch die Zunahme von Gefährdungen bedingte Epidemie führt zu einer Zunahme der Inzidenz.

Eine solche Zunahme der Inzidenz kann verursacht werden

- durch eine Zunahme der exponierten Personen
- durch eine Zunahme der disponierten Personen
- durch eine Zunahme der Wirkdauer einer Exposition (Zeit unter Risiko)
- durch eine Zunahme der force of morbidity

Menschen können gegenüber auf sie wirkende Expositionen immer auch unterschiedlich disponiert, bzw. empfänglich sein. Diese „Empfänglichkeit“ wird oft als eine genetisch-habituelle Eigenschaft vermutet, ist tatsächlich aber auch eine umweltmodifizierbare Eigenschaft. Als solche kann sie sich im Ablauf der epochalen Zeit und in den Generationenfolgen modifizieren, ebenso aber auch innerhalb der Ontogenese und des biografischen Lebensverlaufs. Dispositionen sind, sofern messbar folglich auch epidemiologisch und sozialmedizinisch relevante Klassifikatoren.

Bezeichnet man die Personen, die einem Risiko ausgesetzt sind, als das epidemische Potential und die krankmachenden Einwirkungen als Expositionen, ist eine Epidemie die Folge des Ungleichgewichts von epidemiologischen Potenzialen und der Intensität von Gefährdungen.

Hieraus folgt auch, dass die Struktur der Risikobevölkerung und der Wandel dieser Struktur im Gefolge einer Epidemie deren Verlauf mitbestimmen.

Aus diesen Festlegungen folgt weiter, dass Gefährdungen für die individuelle Gesundheit, bzw. potenzielle Ursachen des individuellen Krankwerdens selbst keine epidemischen Faktoren sind. Epidemische Faktoren sind vielmehr solche, die die Ausbreitung oder die Intensität von Ursachen des Krankwerdens negativ verändern. Sie gehen folglich regelhaft auf soziale, bzw. gesellschaftlich relevante Vorgänge oder Prozesse zurück.

Die Mengenveränderung der Neuerkrankungen während einer Epidemie ist auch das Resultat des sich durch die Epidemie verändernden epidemischen Potentials.

„Niedrige“ Risiken, die auf ein großes Potential einwirken, können quantitativ folgenreicher sein als „hohe“ Risiken, die auf ein kleines Potential wirken. Allein aus der

Mengenveränderung der Zahl neuerkrankender Personen kann also nicht auf die „Gefährlichkeit“ der potenziellen Krankheitsursache geschlossen werden.

Beispiel:

Bestimmte Handlungen, etwa das Rauchen, körperliche Inaktivität oder Fehl- und Überernährung können Krankheiten verursachen. Solche Krankheitsursachen sind jedoch selbst keine epidemischen Faktoren. Epidemische Faktoren sind vielmehr jene, die die Ausbreitung entsprechender Lebensstile in einer Bevölkerung, bzw. in einzelnen ihrer sozialen Schichten bewirken. Das gilt analog für die Zunahme der Intensität solcher Einwirkungen. Analog ist das Zikavirus nicht die Ursache einer Zikaepidemie. Das sind vielmehr jene Ursachen, die die Ausbreitung des Virus und/oder eine Zunahme der „Empfänglichkeit“ für dieses Virus bewirken.

Epidemische Faktoren müssen in der Regel als vom Homo Sapiens selbst verursacht und folglich als soziale Einflüsse identifiziert werden.

Während des Ablaufs einer Epidemie ändern sich immer auch der Umfang und die Struktur des epidemischen Potentials! Es wird gleichsam „verbraucht“.

Während einer Epidemie strebt die Inzidenz zunächst einem Maximum zu, nachfolgend geht sie durch den „Verbrauch“ ihres Potenzials wieder auf ein niedrigeres Niveau zurück. Dies gilt, wenn die entsprechenden zeitlichen Verläufe des epidemischen Zyklus kürzer sind als die Erneuerung des zugehörigen Potenzials.

Sind die Personen einer Bevölkerung gegenüber einer Gefährdung gleich exponiert, jedoch unterschiedlich disponiert, werden sie je nach Disposition zeitlich nacheinander von der Epidemie betroffen werden. Dies gilt auch in der Umkehrung.

Epidemien sind Eigenschaften von Bevölkerung. Es ist die Aufgabe der epidemiologischen Forschung, Erkenntnisse über den Umfang und die Struktur der jeweils Betroffenen/Gefährdeten, also Erkenntnisse über die epidemiologischen Potentiale und ihre Eigenschaften zu gewinnen, um eine Epidemie erfolgreich vermeiden zu können.

9.1.4 Präventive Faktoren

Präventive Faktoren verhindern das Krankwerden oder Verzögern den Beginn einer Erkrankung. Sie sind Ursachen risikobedingter Regressionen.

Präventive Faktoren wirken spiegelbildlich zu den epidemischen. Zuzüglich ist in Rechnung zu stellen, dass auch die Zeitverzögerung des Krankheitsbeginns eine wichtige Präventionswirkung ist. In einem solchen Fall sind die quantitativen Wirkungen denen der Zeitverzögerung von Diagnosestellungen analog.

Da die Disposition gegenüber Krankheit in den einzelnen Altersgruppen unterschiedlich ist, heißt Zeitverzögerung zugleich, den Krankheitsbeginn in einen späteren Lebensabschnitt mit ggf. geringerer/höherer Disposition zu verlagern. Dies macht es hypothetisch denkbar, dass wirksame Prävention die Prävalenz auch vergrößern kann, dies allerdings mit einer „günstigeren“ Altersstruktur der Erkrankten. Die Prüfung solcher Zusammenhänge ist dann die Aufgabe von Analysen zu konkurrierenden Ereignissen.

9.1.5 Interaktion der Bewegungsfaktoren

Die nicht zufälligen Unterschiede in der Verteilung gesundheitlicher Probleme innerhalb einer Bevölkerung lassen sich vollständig auf die epidemiologischen Bewegungsfaktoren zurückführen.

Die dargestellten Bewegungsfaktoren der Prävalenz wirken nie isoliert, sondern in der Regel komplex und wirken in unterschiedlichen Teilbevölkerungen ggf. mit unterschiedlicher Richtung, Intensität und zeitlicher Abfolge.

Diese Festlegung besagt, dass die nichtzufälligen Verteilungsunterschiede Folgen von Bewegungen der Prävalenzen epidemiologischer Fälle sind. Sie sind also auch als komplexe Folgen unterschiedlicher Wirkungen der Bewegungsfaktoren in Teilen einer Bevölkerung zu beschreiben. Solche Unterschiede betreffen die Intensität der wirkenden Bewegungsursachen und die Richtung der Bewegungen.

Hieraus folgt weiter, dass in einer Bevölkerung gleichzeitig jede dieser Bewegungstypen der Prävalenz möglich ist und so zu Verteilungsunterschieden der prävalenten Fälle führt. Der regelhafte Befund der sozialen Ungleichheit der Menschen vor der Belastung mit Krankheit bedarf zur Erklärung folglich einer Abschätzung der Wirkungen der benannten Bewegungsfaktoren jeweils in Teilen der Bevölkerung.