

Workshop

„Real world data“ und Registerdaten in der klinischen und epidemiologischen Forschung: Chancen und Herausforderungen

17./18. November 2016

der Arbeitsgruppen „Statistische Methoden in der Medizin“
(IBS-DR), „Statistische Methoden in der Epidemiologie“ (IBS-DR,
DGEpi), „Statistische Methoden in der klinischen Forschung“
(GMDS), „Epidemiologische Methoden“ (DGEpi, GMDS, DGSMP)

Veranstaltungsort:

Bayer Pharma Aktiengesellschaft
Müllerstraße 178, Gebäude S100, Raum 001
13342 Berlin



Donnerstag, 17. November

13:30-14:00 Begrüßung und Registrierung

Real World Data und Registerdaten in der epidemiologischen Forschung

- 14:00-15:00 Iris Pigeot (Bremen): Nutzung von Sekundärdaten für die pharmakoepidemiologische Forschung: Lassen sich Limitationen der Daten durch moderne statistische Verfahren oder Datenlinkage ausgleichen?
- 15:00-15:20 Daniela Zöller (Mainz): Obtaining a candidate set of covariates for genetic risk scores under data protection constraints in consortia
- 15:20-16:05 Enno Swart (Magdeburg): Gute Praxis Daten-Linkage

Propensity Score Methoden

- 16:25-16:45 Christel Weiss (Heidelberg): Die Anwendung des Propensity Scores bei nicht randomisierten Studien
- 16:45-17:05 Daniela Adolf (Magdeburg): Regressionsanalyse vs. Propensity Score Matching – Methodenvergleich am Beispiel des Herniated-Registers
- 17:05-17:25 Carsten Oliver Schmidt (Greifswald): Propensity-Score Matching im Rahmen der Evaluation einer ambulanten geriatrischen Komplexbehandlung - Work in Progress
- 17:25-17:45 Ingo Langner (Bremen): Reduktion des "Healthy-Sreenee-Bias" in Beobachtungsstudien zum deutschen Mammographie-Screening Programm: ein High-Dimensional-Propensity-Score Ansatz auf Basis von Routinedaten der Krankenkassen

Case studies Real World Data

- 17:55-18:10 Klaus Kaier (Freiburg): Evolution of the volume-outcome relationship among transcatheter aortic valve implantations performed in Germany 2008-2014
- 18:10-18:25 Raphael Peter (Ulm): Adipokines, C-reactive protein and Amyotrophic Lateral Sclerosis – results from a population- based ALS registry in Germany
- 18:25-18:40 Reinhard Meister (Berlin): "Real world data": Die deutsche Embryotox Kohorte - Datenbasis zur Bewertung der Arzneimittelsicherheit in der Schwangerschaft

Gemeinsames Abendessen im Porta Nova (Robert-Koch-Platz 12) ab 19:30 Uhr

Freitag, 18. November

8:30-9:00 Gemeinsame AG-Sitzung (inkl. Sprecherwahlen)

Real World Data und Registerdaten in der klinischen Forschung

- 9:05-10:05 Marc Vandemeulebroecke (Basel): Real-World Evidence in Drug Development
- 10:05-10:25 Christoph Gerlinger (Berlin): Cross Design Synthesis of data from randomized clinical trials and observational studies
- 10:25-10:45 Sebastian Kloss (Berlin): How Real World Evidence can support HTA submissions – Secondary data analyses within the German benefit assessment

Methodische Herausforderungen im Umgang mit Real World Data

- 11:00-11:20 Tobias Bluhmki (Ulm): Methodological Challenges and Time-to-Event Techniques in Diabetes Registry-based Health Services Research
- 11:20-11:40 Regina Stegherr (Ulm): Analysing baseline covariates in studies with delayed entry using a joint model
- 11:40-12:00 Stefan Dahm (Berlin): Trends of cancer survival rates estimated from data of German epidemiologic cancer registries
- 12:00-12:20 Ralph Brinks (Düsseldorf): Berücksichtigung von Kovariablen im Illness Death-Modell
- 12:20-12:40 Annika Hoyer (Düsseldorf): Sequences of prevalence surveys with dependent data – new insights based on spatial statistics

Veranstaltungsabschluss

Nutzung von Sekundärdaten für die pharmakoepidemiologische Forschung: Lassen sich Limitationen der Daten durch moderne statistische Verfahren oder Datenlinkage ausgleichen?

Iris Pigeot, Vanessa Didelez, Dirk Enders, Bianca Kollhorst

Leibniz-Institut für Präventionsforschung und Epidemiologie – BIPS, Achterstraße 30, 28359 Bremen

Im Idealfall werden zur Beantwortung epidemiologischer Fragestellungen Daten im Rahmen von Studien erhoben, die speziell auf die jeweilige Fragestellung zugeschnitten sind. Die qualitätsgesicherte Erhebung von Primärdaten ist aber mit erheblichem Kosten- und Zeitaufwand verbunden und hat zudem mit Schwierigkeiten zu kämpfen, die sich z.B. aus Erinnerungsfehlern der Probanden, einer hohen Non-response oder einer hohen Drop-out Rate ergeben. Daher lohnt es sich u.U., für die Beantwortung bestimmter Fragestellungen bereits existierende Datenquellen, sogenannte Sekundärdaten, zu nutzen. Speziell in der Pharmakoepidemiologie sind zur Untersuchung von seltenen Arzneimittelrisiken oder zum Monitoring eines Arzneimittels nach Marktzulassung sehr große Datensätze erforderlich, um statistisch valide Aussagen treffen zu können, so dass in solchen Fällen häufig auf Sekundärdaten wie z.B. Abrechnungsdaten der Krankenkassen zurückgegriffen wird. Da diese Daten jedoch ursprünglich nicht zu Forschungszwecken erhoben worden sind, sind typischerweise nicht alle notwendigen Informationen in den jeweiligen Datensätzen vorhanden oder sie liegen nur in eingeschränkter Qualität vor.

In diesem Vortrag werden anhand der am BIPS seit 2004 aufgebauten deutschen pharmakoepidemiologischen Forschungsdatenbank GePaRD mit Abrechnungsdaten von vier gesetzlichen Krankenversicherungen mit insgesamt ca. 20 Millionen Versicherten Limitationen dieser Daten und mögliche Ansätze zum Umgang mit diesen Limitationen in praktischen Beispielen aufgezeigt. Dabei werden zum einen statistische Lösungsansätze vorgestellt, aber auch Möglichkeiten diskutiert, fehlende oder nur lückenhaft vorhandene Informationen durch die Verlinkung mit anderen Datenbanken oder mit Primärdaten aufzufüllen. Es wird jedoch nicht nur gezeigt, wie z.B. Primärdaten Lücken in Sekundärdaten füllen können, sondern auch am Beispiel der NAKO Gesundheitsstudie das Potenzial von Sekundärdaten demonstriert, Primärdaten sinnvoll zu ergänzen.

Obtaining a candidate set of covariates for genetic risk scores under data protection constraints in consortia

Daniela Zöller

Institut für Medizinische Biometrie, Epidemiologie und Informatik (IMBEI), Universitätsmedizin der Johannes-Gutenberg-Universität Mainz

In the medical research field, the samples available in a single biobank and corresponding molecular measurements do often not contain enough information to answer increasingly complex scientific questions. By complementing the measurements from one biobank with measurements from other biobanks, statistically more reliable results can be obtained. Yet, due to data protection constraints in Germany, such a contribution of data can only be done at considerable expenses. For example, a written informed consent needs to be obtained of each patient for every scientific question. To ensure that the effort is worthwhile, we propose a method for multivariate regression modeling that requires only aggregated data from different data locations which are permitted to use under the data protection constraints in Germany. This approach is used to identify potentially interesting covariates for a genetic risk score from a potentially large set of covariates.

Specifically, we propose a regularized regression approach, namely componentwise likelihood based boosting, solely based on univariate coefficient estimates obtained from a linear regression for the endpoint of interest and the covariance matrix of the covariates. The univariate coefficient estimates and the covariance matrices of the different data providers can be combined via pooling. Due to regularized regression, data need to be standardized. Ideally, this would be done based on all individual data, but due to the mentioned data protection issues, this might only be possible per data provider. In order to keep the procedure practicable, we propose to use a heuristic and a block-heuristic version of the above given method. These versions can reduce the amount of data that needs to be transferred between the participating locations and the number of needed calls for data considerably.

We evaluate our proposed method in a simulation study. Here, we compare our results to the results obtained with the pooled individual data. The proposed method is seen to lead to a moderate increase in the number of false positive detected covariates, but is able to detect covariates with true effect to a comparable extent as a method based on the individual data, if the number of locations is rather small. With increasing number of locations, the number of false positives increases considerably. The influence of standardizing per data provider instead of standardizing the total data is seen to be even favorable in our simulations. In summary, the proposed approach is seen to provide multivariable risk prediction models with an adequate performance in situations with distributed data.

Gute Praxis Datenlinkage

Enno Swart¹, Stefanie March¹, Manfred Antoni², Joachim Kieschke³, Bianca Kollhorst⁴, Birga Maier⁵, Gabriele Müller⁶, Murat Sariyar⁷, Mandy Schulz⁸, Jan Zeidler⁹, Falk Hoffmann¹⁰ für die Arbeitsgruppe Erhebung und Nutzung von Sekundärdaten (AGENS) der Deutschen Gesellschaft für Sozialmedizin und Prävention (DGSMP) und der Deutschen Gesellschaft für Epidemiologie (DGEpi) sowie für die Arbeitsgruppe Validierung und Linkage von Sekundärdaten des Deutschen Netzwerks für Versorgungsforschung (DNVF)

¹ Institut für Sozialmedizin und Gesundheitsökonomie (ISMG), Medizinische Fakultät, Otto-von-Guericke-Universität Magdeburg, Magdeburg

² Institut für Arbeitsmarkt- und Berufsforschung (IAB) der Bundesagentur für Arbeit, Nürnberg

³ Epidemiologisches Krebsregister Niedersachsen, Registerstelle, Oldenburg

⁴ Abteilung Biometrie und EDV, Leibniz-Institut für Präventionsforschung und Epidemiologie – BIPS, Bremen

⁵ Berliner Herzinfarktregister am Fachgebiet Management im Gesundheitswesen an TU Berlin, Berlin

⁶ Zentrum für Evidenzbasierte Gesundheitsversorgung (ZEGV), Universitätsklinikum und Medizinische Fakultät Carl Gustav Carus, TU Dresden, Dresden

⁷ TMF – Technologie- und Methodenplattform für die vernetzte medizinische Forschung e.V., Berlin

⁸ Zentralinstitut für die kassenärztliche Versorgung (ZI), Berlin

⁹ Center for Health Economics Research Hannover (CHERH), Leibniz Universität Hannover, Hannover

¹⁰ Department für Versorgungsforschung, Fakultät für Medizin und Gesundheitswissenschaft, Carl von Ossietzky Universität Oldenburg, Oldenburg

Die Verknüpfung verschiedener Datenquellen, genannt Datenlinkage oder auch Record Linkage, zur Beantwortung komplexer Fragestellungen findet in den letzten Jahren in Deutschland vermehrt Anwendung. Jedoch mangelt es bisher an publizierten Erfahrungen. Neue Projekte erarbeiten sich in der Regel autark voneinander das notwendige Handwerkszeug. Daher hat sich eine Gruppe von Forschern zusammengefunden, um ihre Erfahrungen im Rahmen einer Guten Praxis Datenlinkage als mögliche Hilfestellung für Projekte, Gutachter sowie Datenschützer und Ethikkommissionen zusammenzustellen. Ziel dieser Guten Praxis Datenlinkage ist es deshalb, eine Unterstützung für zukünftige Projekte zu liefern, die Daten aus Deutschland auf individueller Ebene verknüpfen möchten. Neben den (datenschutz-)rechtlichen Rahmenbedingungen werden dabei auch praxisorientiert die Arten des Datenlinkage, deren Anwendungsfelder und mögliche Fallstricke sowie Ansätze zu deren Vermeidung anhand von Beispielen dargestellt. Abschließend sollen (langfristige) Visionen eine Diskussion anregen, die das Datenlinkage auch in Deutschland erleichtern könnte und die positive Entwicklung weiter vorantreibt, auch im Hinblick auf eine internationale Anschlussfähigkeit.

Die Anwendung des Propensity Scores bei nicht randomisierten Studien

Christel Weiß, Dennis Ferdinand

Randomisierte Studien gelten als Goldstandard der klinischen Forschung. Aufgrund ethischer, finanzieller und zeitlicher Limitationen ist die Durchführung von randomisierten klinischen Studien jedoch nicht immer möglich. In diesem Fall erfolgt die Zuteilung der Teilnehmer zur jeweiligen Therapiegruppe aufgrund festgelegter klinischer Kriterien oder aber auf Basis persönlicher Entscheidungen der Studienleitung. Es erscheint naheliegend, dass hierbei die Möglichkeit besteht, dass sich die beiden Gruppen zu Beginn der Studie bezüglich eines oder mehrerer Merkmale signifikant unterscheiden. Die Folge dessen zeigt sich in der Interpretation der Studienergebnisse, die nicht durch Therapieunterschiede bedingt sind, sondern auch im ungleichen Ausgangszustand der Teilnehmer zu suchen sind. Um diesem Ungleichgewicht entgegenzuwirken, nutzt der Propensity Score die Summe aller dem Untersucher bekannten Merkmale eines Teilnehmers und subsummiert sie in einem einzigen Vektor [1, 4].

Die Anwendung dieses Scores wird am Beispiel einer prospektiven, nicht-randomisierten Studie der Adipositaschirurgie der Universitätsmedizin Mannheim dargelegt. In dieser Studie werden die Auswirkungen der laparoskopischen OP-Verfahren des Schlauchmagens und des Roux-Y-Magenbypasses hinsichtlich ihrer Effekte auf den postoperativen Gewichtsverlust sowie die Körperzusammensetzung, also die Veränderung von Zellmasse, Körperfett und Körperwasser verglichen [3]. Es wird gezeigt, wie der Score berechnet wird und wie er zu einem unverzerrten und objektiven Therapievergleich beitragen kann [2].

- [1] Austin PC (2011): An Introduction to Propensity Score Methods for Reducing the Effects of Confounding in Observational Studies. *Multivariate behavioral research* 46, 399-424.
- [2] Ferdinand D, Weiss C (2016): Get the most from your data: a propensity score model comparison on real-life data. *Int J Gen Med*, 9, 123-131
- [3] Otto M., Elrefai M, Krammer J, Weiß C, Kienle, P, Hasenberg T (2015): Sleeve Gastrectomy and Roux-en-Y Gastric Bypass Lead to Comparable Changes in Body Composition after Adjustment for Initial Body Mass Index. *Obesity surgery*, 1-7.
- [4] Rosenbaum PR, Rubin DB (1983): The centrale role of the propensity score in observational studies for causal effects. *Biometrika* 70, 41-55.

Regressionsanalyse vs. Propensity Score Matching – Methodenvergleich am Beispiel des Herniamed-Registers

D. Adolf, G. Kaldasheva, M. Hukauf, T. Keller, F. Köckerling

Hintergrund: Herniamed ist ein multizentrisches, internetbasiertes Hernienregister, in das derzeit mehr als 500 Kliniken und niedergelassene Chirurgen in Deutschland, Österreich, der Schweiz und Italien ihre Operationen eingeben. Die Ergebnisse der Behandlungen werden bis zu zehn Jahre nachverfolgt.

Methoden: Ziel einer statistischen Auswertung dieser Daten ist oftmals die Untersuchung des Einflusses spezieller Risikofaktoren oder der Vergleich verschiedener Operationsmethoden. Hierbei müssen jedoch jeweils weitere Einflussfaktoren berücksichtigt werden. Um dies zu realisieren eignen sich einerseits multivariable Modelle, speziell Regressionsverfahren. Aussagen über einzelne Einflussparameter sind dann möglich, da gleichzeitig für die Effekte der anderen Größen adjustiert wird. Andererseits können univariate Vergleiche an Stichproben durchgeführt werden, die über Matching-Verfahren aufeinander abgestimmt und somit auch für die dabei verwendeten Parameter adjustiert sind. In unserem Beispiel werden logistische Regressionsmodelle mit den Ergebnissen nach Propensity Score Matching gegenüber gestellt, wobei die weiteren potentiellen Einflussparameter im Regressionsmodell bei Bildung der Scores als Matchingparameter dienen. Das Propensity Score Matching wird über einen Greedy-Algorithmus mit Calipern von 0.1 bis 1.25 Standardabweichungen realisiert.

Ergebnisse: Grundlage des Vergleichs der beiden statistischen Vorgehensweisen sind Daten zur Analyse mehrerer Operationsverfahren bei Leistenhernien. Im Ergebnis wird zunächst die Güte des Matchings als auch die Anzahl der einander zugeordneten Paare in Abhängigkeit vom Caliper betrachtet. Weiterhin werden die Odds Ratio-Schätzer des interessierenden Einflussparameters aus Regression und Matching gegenüber gestellt.

Propensity-Score Matching im Rahmen der Evaluation einer ambulanten geriatrischen Komplexbehandlung - Work in Progress

Schmidt, Carsten Oliver, Kiel, Simone; Zimak, Carolin; Chenot, Jean-François

Universitätsmedizin Greifswald, Fleischmannstr. 8, 17475 Greifswald

Hintergrund

Die ambulante geriatrische Komplexbehandlung (AGKB) ist ein Programm der AOK Nordost zur Integrierten Versorgung (§140a SGB V) für Patienten mit geriatritypischer Multimorbidität, um Hospitalisierungen, Stürze und Pflegebedürftigkeit zu vermeiden oder zu verzögern.

Ergebnisse aus der Versorgungsforschung zur Effektivität und Nachhaltigkeit liegen dazu bislang nicht vor. Mit einem Propensity Score Matching unter Nutzung von Routinedaten der AOK Nordost werden Effektivität und Nachhaltigkeit der AGKB hinsichtlich der Endpunkte Pflegestufe, Heimaufnahme, Krankenhausaufnahme und Kosten evaluiert.

Dieser Beitrag diskutiert die Methodik und Problematiken zur Erzielung eines möglichst gutes Matchings. Eine besondere Herausforderung dabei war die Selektion von Kontrollen mit geeignetem Bezug zum Interventionsquartal der Fälle.

Methoden

Für das Matching wurden Routinedaten von AGKB Teilnehmern und Kontrollen verwendet, die danach ausgewählt wurden, sowohl einen Einfluss auf die Teilnahme an der AGKB als auch auf die Endpunkte der Intervention zu haben. Die vier der AGKB-Versorgung vorausgegangenen Quartale sowie das Indexquartal bilden den Beobachtungszeitraum für die Variablen, die in das Matching einbezogen wurden. In einem weiteren Szenario wurde das Indexquartal ausgeschlossen.

Das Matching wurde in den folgenden Schritten realisiert:

1. Im ersten Schritt fand anhand von Alter, Geschlecht, Behandlungsindikationen, und ausgewählten Versorgungsparametern ein Prämatching mittels eines exakten Matchings mit >100 Kontrollen pro Fall statt, um einen korrekten Bezug der Kontrollen auf das Indexquartal der jeweils gematchten AGKB Fälle zu ermöglichen. Kontrollen wurden zwischen Quartalen mit Wiederholung gezogen.
2. Im zweiten Schritt wurden Propensity Scores mittels logistischer Regression unter Nutzung des vollen Spektrums von Prädiktoren berechnet.

3. Im dritten Schritt fand die eigentliche Zuweisung der Kontrollen quartalsweise mit Ziehung ohne Wiederholung statt. Jedem Fall wurden vier Kontrollen zugeordnet.

Die Qualität des Matchings wurde u.a. mit der Standardized Mean Difference (SMD) bewertet.

Ergebnisse

699 Patienten mit einem Durchschnittsalter von 79 Jahren erhielten im Zeitraum 2009 bis 2013 an drei Standorten in Mecklenburg-Vorpommern eine ca. 20-tägige AGKB. Indexquartale verteilten sich auf 19 Quartale. Zum Matching standen 251.000 Kontrollen der AOK Nordost zur Verfügung. Insgesamt 67 Fälle wurden wegen seltener Erkrankungen bzw. extremen Versorgungskosten ausgeschlossen. Auf Basis des Prämatchings wurde ein Pool von ca. 87.000 Kontrollen (z.T. mit Wiederholung) ausgewählt. Die SMD der einzelnen Prädiktoren war immer <10% und im Median <3%. Über alle Matchings betrug die mediane SMD 1,4% mit einem Range von -4.4 – 4.9.

Diskussion

Auch bei Verfügbarkeit einer großen Kontrollpopulation ist die geeignete Zuweisung von Kontrollen komplex. Mittels exakten Matchings sind nur bei Berücksichtigung weniger Variablen geeignete Kontrollen zu finden. Die Qualität des Matchings ist hinsichtlich der SMD als zufriedenstellend zu bezeichnen. Bei unserem Vorgehen weisen Kontrollen korrekte Bezüge zu den Beobachtungszeiträumen der Fälle auf. Die Verfügbarkeit mehrerer Matchings mit ähnlicher Qualität erleichtert aussagekräftige Sensitivitätsanalysen. Auswertungen sind derzeit noch nicht abgeschlossen.

Reduktion des "Healthy-Sreenee-Bias" in Beobachtungsstudien zum deutschen Mammographie-Screening-Programm: ein High-Dimensional-Propensity-Score Ansatz auf Basis von Routinedaten der Krankenkassen

Langner I, Zeeb H, Enders D, Giersiepen K, Hense HW

Hintergrund: Im deutschen Mammographie-Screening-Programm (MSP) werden Frauen ohne prävalenten Brustkrebs (BK) im Alter zwischen 50 und 69 Jahren im Rhythmus von 2 Jahren mit einem zeitnahen Termin zur Mammographie eingeladen. Das MSP startete 2005 unter der gesetzlichen Vorgabe, die Auswirkungen des MSP auf die Brustkrebsmortalität zu evaluieren. Allerdings wurden begleitend zum MSP keine Daten gesammelt, die eine direkte Evaluation der Effekte des MSP zulassen würden. Daher sollte im Rahmen einer Machbarkeitsstudie geprüft werden, ob andere Sekundärdatenquellen für eine Evaluation genutzt werden können. Dabei besteht ein wichtiges Kernproblem darin, dass rückwirkend nicht bestimmt werden kann, wann eine Frau sich gegen die MSP-Teilnahme entschieden hat: in den Datenquellen können nämlich nur Teilnehmerinnen (TN+) mit dem Termin der MSP-Untersuchung identifiziert werden, eine entsprechender Einladungstermin für Nichtteilnehmerinnen (TN-) liegt dagegen nicht vor.

Fragestellung: Beobachtungsstudien zur Evaluation des deutschen Mammographie-Screening-Programm (MSP) unterliegen Verzerrungen, die auf der individuellen Teilnahmebereitschaft bzw. –fähigkeit der Zielpopulation des MSP beruhen. Auf der Basis von Krankenkassendaten, die umfangreiche Informationen zu potentiellen, gesundheitsrelevanten Confoundern enthalten, haben wir untersucht, ob High-Dimensional-Propensity-Score (HDPS) Verfahren zur Reduktion der Verzerrungen genutzt werden können.

Methoden: Datengrundlage bildeten die Abrechnungsdaten mehrerer Krankenkassen des Zeitraums 2004 bis 2012, die pro Datenjahr etwa 3,5 Millionen Frauen im Alter 50 bis 69 umfassten und für eine Teilpopulation mit amtlichen Todesursachen angereichert wurden. Wir haben explorativ hochdimensionale Propensity-Modelle (PM) untersucht, die eine Balancierung der Studienpopulation mittels HDPS-Matching ermöglichen und so einen in den Rohdaten erkennbaren Healthy Sreenee Bias (HSB) in den Analysen ausgleichen sollten. Die PM enthielten Faktoren aus den patienten-individuellen Datendimensionen Diagnosen, Operationen und Prozeduren, ambulante Behandlungen, Teilnahme an anderen Vorsorgemaßnahmen, Häufigkeit von Krankenhausaufenthalten und Arztbesuchen, sowie regionale Faktoren wie Urbanität, Deprivation und Arztdichten. Der MSP-Teilnahmestatus sowie PM-Faktoren wurden für einen 2-jährigen Baseline-Zeitraum bei bis dahin BK-freien Frauen (Ausschluss prävalenter BK) erhoben. Den TN- wurden Pseudo-Einladungsdaten zugeordnet, die in der zeitlichen Verteilung den Untersuchungsdaten der TN+ entsprachen. Der MSP-Teilnahmestatus (Teilnahme versus Nicht-Teilnahme) wurde altersadjustiert mittels logistischer Regression für die Endpunkte kumulative Gesamtmortalität (kGM) sowie Inzidenz-basierte Brustkrebsmortalität (kIBM) innerhalb von 4 Jahren nach dem Baseline-Zeitraum analysiert.

Das HDPS-Matching und die Regressionsanalyse wurden je Endpunkt einhundertmal wiederholt, um zufällige Effekte des Matchings auszugleichen; die präsentierten Odds-Ratio-Schätzer sowie die zugehörigen 95% Konfidenzintervallgrenzen (KI) wurden als Mittelwerte der Ergebnisse dieser einhundert Analysen ermittelt. Unter der Annahme, dass die MSP-Teilnahme noch keinen Effekt auf die kurzfristige kGM und kIBM haben kann, wurde die relative Bias-Reduktion (in %) aus dem

Vergleich der Odds-Ratio-Schätzer (OR) ermittelt, die sich bei den Analysen mit und ohne HDPS-Matching ergaben.

Ergebnisse: In das finale PM wurden 542 Faktoren einbezogen. Zusätzliche Faktoren brachten keine weitere Verbesserung in der Balancierung der HDPS-gematchten Studienpopulation. In der ungematchten Population zeigte sich bei der kGM für den MSP-Einfluss ein OR von 0.506 (95%KI 0.474-0.540) und bei der kIBM ein OR von 0.397 (95%KI 0.276-0.571) gegenüber einem OR von 0.778 (95%KI 0.715-0.847) (kGM) und 0.711 (95%KI 0.440-1.148) (kIBM) in der PS-gematchten Population. Die Einführung eines Indexdatums bei TN- gemäß der altersspezifischen Verteilung der MSP-Teilnahmetermine der TN+ und ein Ausschluss der TN- mit einem inzidenten Brustkrebs vor dem Indexdatum in der ungematchten Population führte zu einem OR des MSP-Einflusses bei kIBM von 0.601 (95%KI 0.410-0.880). Eine auf dieser Population basierende, HDPS-gematchte Population ergab in der Analyse ein OR von 0.957 (95%KI 0.563-1.627) für den MSP-Einfluss bei kIBM. Im Vergleich zum OR der ungematchten Population konnte der Abstand des Punktschätzers zu einem OR=1 (=angenommener, Bias-freier, wahrer Effekt für MSP-Teilnahme auf kurzfristige, kumulative Mortalität) bei der kGM um 58%, bei der kIBM um 89% reduziert werden.

Diskussion: Die Evaluation des MSP hinsichtlich der Auswirkungen auf die Brustkrebsmortalität auf Basis von GKV-Daten bietet ein Paradebeispiel für die Verwendung eines HDPS-Matchings zur Balancierung von Verzerrungen in der Studienpopulation. Zum einen sind die beim Einfluss der MSP-Teilnahme hinsichtlich der Mortalität wirkenden Confounder unbekannt, zum anderen bietet der Studienansatz auf Basis von Kassendaten die Möglichkeit, auf individueller Ebene neben der MSP-Teilnahme auch eine Vielzahl von gesundheitsrelevanten Faktoren mit einzubeziehen. Weiterhin bietet die Teilnahmerate von ca. 50% bis 60% eine ausreichende Grundlage, um in einem PM mit vielen Einflussfaktoren die notwendige Power zu erreichen, während das Zielereignis "Tod durch Brustkrebs" ein eher seltenes Ereignis ist, so dass sich die Einbindung vieler einzelner Confounder in ein konventionelles Regressionsmodell ausschließt. Da die anspruchsberechtigten Frauen nicht am Einladungsdatum als TN+ und TN- identifiziert werden können, musste zusätzlich zu einer Balancierung mittels HDPS-Matchings auch die Einladungszeit (als Proxy für „exposure time zero“) modelliert werden. Schließlich wurde das Eingangskriterium „BK-frei“ durch die ausschließliche Benutzung der inzidenz-basierten Mortalität auch bei TN- erfüllt. Mit einem entsprechend angepassten Design konnte durch die Balancierung der Studienpopulation für die kIBM eine Bias-Reduktion von insgesamt 93% erreicht werden.

Evolution of the volume-outcome relationship among transcatheter aortic valve implantations performed in Germany 2008-2014

Klaus Kaier¹, Werner Vach¹ and Jochen Reinöhl²

¹ Institute of Medical Biometry and Statistics, Faculty of Medicine and Medical Center – University of Freiburg, Germany

² Department of Cardiology and Angiology I, Heart Center Freiburg University, Germany

Abstract

We examine the volume–outcome relationship with respect to isolated transcatheter aortic valve implantations (TAVI) using patient characteristic and in-hospital outcome data from the DRG Statistics of all TAVI procedures performed in Germany between its introduction in 2008 and 2014 (N=43,996). Our question was whether the volume-outcome relationship for TAVI exists on the center level, whether it occurs equally for different outcomes, and how it develops over time.

Logistic and linear regression analyses were carried out for the endpoints in-hospital mortality, bleeding, stroke, probability and length of ventilation, length of hospital stay, and reimbursement. Risk-adjustment was applied using a predefined set of patient characteristics¹ to account for differences in the risk factor composition of the patient populations between centers and over time.

We found that risk-adjusted in-hospital mortality steadily decreases over the years and is lower the higher the annual procedure volume at the respective center is (<50, 50-99 and ≥ 100 procedures). The magnitude of this effect declines over the observation period. Regarding all possible endpoints, our results indicate some sort ceiling effect of the volume-outcome relationship: The volume-outcome relationship is eminent in circumstances of relatively unfavorable outcomes (e.g. during early years after TAVI introduction in Germany and with respect to relatively high mortality, bleeding, and ventilation rates). Alongside improving outcomes (e.g. in later years and with respect to the same outcomes), however, the volume-outcome relationship decreases. Interestingly, a volume-outcome relationship seems to be absent in circumstances of constantly low outcome rates (e.g. with respect to stroke rates which are constantly below 2 %).

We conclude that the hypothesized volume-outcome relationship for TAVI exists but diminishes in accordance to a reduced capacity for improvement.

References

1. Reinöhl, J., K. Kaier, H. Reinecke, C. Schmoor, L. Frankenstein, W. Vach, A. Cribier, F. Beyersdorf, C. Bode and M. Zehender. Effect of Availability of Transcatheter Aortic Valve Replacement on Clinical Practice. *N. Engl. J. Med.* **373**, 2438–2447 (2015).

Adipokines, C-reactive protein and Amyotrophic Lateral Sclerosis – results from a population- based ALS registry in Germany

Raphael S Peter¹, Gabriele Nagel¹, Angela Rosenbohm², Wolfgang Koenig^{3,4,5}, Luc Dupuis⁶, Dietrich Rothenbacher¹, Albert C Ludolph² and the ALS Registry Study Group

¹ Institute of Epidemiology and Medical Biometry, Ulm University, Ulm, Germany

² Department of Neurology, Ulm University, Ulm, Germany

³ Department of Internal Medicine II - Cardiology, University of Ulm Medical Centre, Ulm, Germany

⁴ Deutsches Herzzentrum München, Technische Universität München, Munich, Germany

⁵ DZHK (German Centre for Cardiovascular Research), partner site Munich Heart Alliance, Munich, Germany

⁶ INSERM U1118, Université de Strasbourg, Strasbourg, France

Abstract

Between October 2010 and June 2014 a population-based case-control study has been established based on the ALS registry Swabia, Southern Germany. Data of 289 ALS patients (mean age of 65.7 (SD 10.5) years, 59.5% men) and 506 controls were included. During median follow-up of 14.5 months of 279 ALS patients 104 (53.9 % men, 68.9 (10.3) years) died. Blood samples were measured for leptin, adiponectin and hs-CRP. Information on sociodemographic and lifestyle factors was assessed by a standardized questionnaire. Conditional logistic regression was used to calculate multivariate odds ratios (OR) for risk of ALS. Survival models were used in cases only to appraise the prognostic value.

Compared to controls, ALS patients were characterized by lower levels of school education, BMI and smoking prevalence. Adjusted for covariates, leptin was inversely associated with risk of ALS (top vs. bottom quartile: OR 0.49; 95%CI 0.29-0.80), while for adiponectin a positive association was found (OR 2.89; 95%CI 1.78-4.68, top vs. bottom quartile). Among ALS patients increasing leptin concentrations were associated with longer survival (p for trend 0.002), while for adiponectin no association was found (p for trend 0.55). In contrast, hs-CRP was neither associated with onset nor prognosis of ALS.

Leptin and adiponectin, two key hormones regulating energy metabolism, were strongly and independently related with risk of ALS. Leptin levels were further negatively related with overall survival of ALS patients.

**„Real world data“:
Die deutsche Embryotox Kohorte – Datenbasis zur Bewertung der
Arzneimittelsicherheit in der Schwangerschaft**

Reinhard Meister

Beuth Hochschule für Technik, Berlin, Fachbereich Mathematik, Physik, Chemie

Informationen zur Arzneimittelsicherheit in der Schwangerschaft helfen in erster Linie dabei, unnötige Ängste bei medizinisch notwendigem Medikamentengebrauch zu vermeiden. Natürlich sind Warnhinweise und Handlungsempfehlungen bei tatsächlich erhöhtem Risiko angebracht.

Da es bei Schwangeren keine kontrollierten Studien gibt und in Deutschland keine umfangreichen Registerdaten zu Arzneimitteln und Schwangerschaftsverlauf zur Verfügung stehen, sind die Daten des Berliner Pharmakovigilanz- und Beratungszentrums für Embryonaltoxikologie die wichtigste Basis für empirische Studien mit nationalen Daten.

Die Daten enthalten, prospektiv erhoben, eine systematische Anamnese auch vorheriger Schwangerschaften, Angaben zu Erkrankungen und Arzneimittelexpositionen in der Schwangerschaft sowie Informationen zum Ausgang der Schwangerschaft, die etwa 8 Wochen nach dem errechneten Geburtstermin erfragt werden.

Wie aus der Idee, kostenlose Beratung mit der Bitte um Zustimmung zur Dokumentation und wissenschaftlichen Verwertung von Daten der Anfragenden zu verbinden, eine wertvolle Datensammlung von mittlerweile mehr als 45000 prospektiv erfassten Schwangerschaftsverläufen entstehen konnte, ist Inhalt des Beitrags.

Struktur, Prinzipien der Dokumentation, Entwicklung der Kohorte sowie statistische Aspekte und Ergebnisse ausgewählter Studien werden vorgestellt und erläutert.

Literaturhinweise

www.embryotox.de

Schaefer C, Ornoy A, Clementi M, Meister R, Weber-Schoendorfer C (2008) Using observational cohort data for studying drug effects on pregnancy outcome—methodological considerations. *Reproductive Toxicology* 26 (1), 36-41

Real-World Evidence in Drug Development

Marc Vandemeulebroecke, Novartis Pharma AG

In this introductory talk, I explore the concept of “real-world evidence” and related terms: what are the perspectives and motivations of different stakeholders in drug development? What drives the gap between efficacy and effectiveness? What questions are addressed by clinical trials and real-world studies, and how do their designs differ? What are the consequences of these concepts for the process of approval and market access for new medicinal products? I discuss opportunities and challenges of real-world evidence in drug development. In conclusion, the role of real-world data and evidence can be expected to grow. However, it is important to be clear about their purpose, value and limitations, also in contrast to clinical trials.

Cross Design Synthesis of data from randomized clinical trials and observational studies

PD Dr. Christoph Gerlinger* Dr. Tatsiana Vaitsiakhovich[†]
Dr. Anna Filonenko[‡]

Abstract

Data on the benefits and risks of drug arise from multiple sources like randomized clinical trials, observational studies, or billing data from health insurers. Typically the data from the different sources are analyzed and synthesized separately.

Randomized clinical trials include a highly selected subset of the total patient population in order to have the highest possible internal validity. As a consequence their external validity is reduced. Results of observational studies may be more representative for routine clinical practice with focus on external validity. However their internal validity is compromised due to confounding and selection bias. It remains unclear whether the different treatment effects are due to the different inclusion criteria or are caused by other "real life" factors.

There is growing need to maximize internal and external validity of the treatment efficacy, safety, adherence, health related quality of life, and other outcomes collected in randomized clinical trials and observational studies to inform healthcare access and policy. The methods to synthesize data from randomized clinical trials and from observational studies into a single analysis have received research attention recently in the biomedical and biostatistical literature.

We will present here the cross design synthesis method that was initially recommended by the United States General Accounting Office as a new strategy for medical effectiveness research. We will also present a worked example from long acting reversible contraceptives.

*Research & Development Statistics, Bayer Pharma AG, 13353 Berlin, Germany and Gynecology, Obstetrics and Reproductive Medicine, University Medical School of Saarland, 66421 Homburg/Saar, Germany

[†]RLE Strategy & Outcomes Data Generation, Bayer Pharma AG, 13353 Berlin, Germany

[‡]Market Access Pulmonology / Innovative Women's Healthcare, Bayer Pharma AG, 13353 Berlin, Germany

How Real World Evidence can support HTA submissions – Secondary data analyses within the German benefit assessment

S. Kloss¹, E. Basic², F. Leverkus²

¹Bayer Business Services

²Pfizer Deutschland GmbH

Background

“Real World Evidence” is getting increasingly important and is broadly seen complementary to randomized controlled trials (RCTs). Data collection and data analyses within observational research using existing data sources can be done time- and cost-efficiently. In Germany, IQWiG and G-BA do not consider results from observational studies to assess the additional benefit of a new medication. However, these types of studies can still be used to support HTA submissions in terms of real world incidence and prevalence estimates to understand the medical need and impact, costs of therapies, standard of care evaluation or in order to show transferability of RCT findings to the German healthcare system.

Methods

In order to estimate the number of patients who might benefit from a treatment for venous thromboembolism (VTE, i.e. deep vein thrombosis or pulmonary embolism), three database studies were set up using different data sources for estimating the prevalence and incidence of VTE in Germany. Results were intended to be included in a benefit dossier for a new indication of an anticoagulation therapy. The data sources were chosen as: IMS Disease Analyzer® (primary care physician panel), Elsevier HRI database (claims database) as well as the DIMDI database (Statutory health insurances in Germany).

Results

Results for prevalence and incidence of VTE in Germany differed between the data sources. Results of the DIMDI database and HRI database looked similar although time periods were different. There was a strong overestimation regarding the number of patients with diagnoses in the IMS Disease Analyzer®. IQWiG acknowledged results from these studies in terms of transferability from RCTs to the German healthcare system. Methods, results as well as strengths and limitations were discussed in the oral hearing session with the G-BA.

Discussion

Different data sources come along with different strengths and limitations. It is critical to understand which research question can be addressed and answered by which data source, knowing very well the context of the data, patient characteristics as well as how the data was captured. Where claims data might be understood as well structured and easy to use, a lot of information is lacking (e.g. laboratory values). Data from primary care physician panels could add additional evidence but should be handled with caution as interpretation is often not straightforward. Generally, it is challenging to compare data sources from different healthcare sectors. Therefore, a detailed description of the data source for all its limitations as well as including the study protocol with the underlying methods is crucial in a potential HTA submission.

Methodological Challenges and Time-to-Event Techniques in Diabetes Registry-based Health Services Research

Tobias Bluhmki¹, Peter Bramlage², and Jan Beyersmann¹

¹Institute of Statistics, Ulm University, Helmholtzstrasse 20, 89081 Ulm, Germany
email: tobias.bluhmki@uni-ulm.de, jan.beyersmann@uni-ulm.de

²Institute of Pharmacology and Preventive Medicine, Menzelstrasse 21, 15831 Mahlow, Germany
email: peter.bramlage@ippmed.de

Objective: Complex longitudinal sampling and the observational structure of patient registries in health services research are associated with methodological challenges. Based on an ongoing diabetes registry, we exemplify common pitfalls and want to stimulate discussions on the design, development, and deployment of future longitudinal patient registries and registry-based studies.

Study Design and Setting: We employ data from the prospective, observational German Diabetes Versorgungs-Evaluation (DIVE) registry [1]. A principle aim of the registry was to explore predictors for the initiation of a basal insulin supported therapy in patients with type-2 diabetes initially prescribed to glucose-lowering drugs alone [2].

Results: Major challenges are the adequate choice of the timescale of interest, missing mortality as well as baseline information, time-dependent outcomes, delayed study entries, different follow-up times, and competing events. We show that time-to-event methodology is a valuable tool for improved statistical evaluation and should be preferred over simple case-control approaches.

Conclusion: Patient registries provide rich data sources for health services research. Analyses are subject to a trade-off between data availability, clinical plausibility, and statistical feasibility. Cox' proportional hazards model allows for the evaluation of the outcome-specific hazards, but the current setting only allows an investigation of covariates measured at (delayed) study entry [3]. Further, personalized prediction of patient outcomes is compromised by missing mortality information.

- [1] Danne T, Kaltheuner M, Koch A, Ernst S, Rathmann W, Russmann HJ, et al. 'Diabetes Versorgungs-Evaluation' (DIVE) – a national quality assurance initiative at physicians providing care for patients with diabetes. *Deutsche Medizinische Wochenschrift* 2013; 138(18):934-9
- [2] Danne T, Bluhmki T, Seufert J, Kaltheuner M, Rathmann W, Beyersmann J, et al. Treatment intensification using long-acting insulin – predictors of future basal insulin supported oral therapy in the DIVE registry. *BMC Endocrine Disorders*. 2015;15(1):1.
- [3] Keiding N, Knuiman M. Letter to the editor survival analysis in natural history studies of disease. *Stats Med* 1990; 9(10):1221-2.

Analysing baseline covariates in studies with delayed entry using a joint model

Regina Stegherr, Tobias Bluhmki, Jan Beyersmann

Institute of Statistics, Ulm University, Helmholtzstrasse 20, 89081 Ulm, Germany

The time of study entry into a randomized clinical trial is a natural choice for "time zero", but in other life history cohorts study entry may happen after time origin, leading to left-truncated data. E.g., in observational health services data on diabetes patients a possible time origin would be diagnosis of diabetes. Some patients may enter the study upon diagnosis, but others may have a known date of diagnosis before start of data collection. A relevant baseline covariate would be Body Mass Index (BMI), the problem being that such data may be measured upon study entry and, hence, not at baseline for those with a delayed study entry. The problem has been succinctly summarized in a letter by N Keiding and M Knuiman, *Statistics in Medicine*, Vol. 9, 1221-1222 (1990). The aim of the present work is to investigate whether a joint model for a longitudinal covariate such as BMI and time to disease event may be used to analyse the impact of baseline covariates, possibly unmeasured due to delayed study entry, on the time to event. This is in contrast to standard use of a joint model where the typical aim is to investigate the impact of the current value of the longitudinal marker, possibly unmeasured because of non-continuous updating, on the hazard of an event.

Trends of cancer survival rates estimated from data of German epidemiologic cancer registries

Stefan Dahm, Zentrum für Krebsregisterdaten, Berlin

Objectives:

Cancer survival is a key measure in evaluation of cancer control. Large temporal and geographical variation in survival rates estimated from German epidemiologic cancer registries make it difficult to interpret trends in survival.

Methods:

The study is based on 4.0 million cancer cases notified in 14 German epidemiologic cancer registries during the years from 1997 to 2013. 1- and 5-year relative survival rates were estimated stratified by cancer site, registry, follow-up year, age group and sex. For each stratum the proportion of death certificate only (DCO) notifications was recorded in addition. Survival rates by cancer sites, age groups and sex were modeled by mixed linear regression using random intercepts for registries and DCO proportion and follow-up year as fixed effects. Based on this linear model annual cancer survival rates within Germany were predicted at DCO proportions of 0 %.

Results:

After controlling for DCO proportions and follow-up years survival rates showed only a small variation between registries. For most cancer sites we estimated significant positive trends for survival rates over the years. The strongest increase of survival was seen for prostate cancer

Conclusions:

The proposed method could be used to estimate recent trends in cancer survival in Germany.

Berücksichtigung von Kovariablen im Illness-Death-Modell

Ralph Brinks^{1,2}, Sophie Kaufmann², Annika Hoyer²

1) Hiller Forschungsinstitut für Rheumatologie, Universitätsklinikum Düsseldorf

2) Institut für Biometrie und Epidemiologie, Deutsches Diabetes-Zentrum (DDZ), Leibniz-Zentrum für Diabetes Forschung an der Heinrich-Heine-Universität Düsseldorf,

In [1] wurde basierend auf dem Illness-Death-Modell (IDM) eine partielle Differentialgleichung vorgestellt, die die Prävalenz einer Erkrankung mit den Übergangsraten (Inzidenz- und Mortalitätsraten) des IDM in Relation setzt.

Für die Abschätzung des Effektes von populationsbezogenen Interventionen ist es wichtig, die Wirkung von Kovariablen im IDM beschreiben zu können. In diesem Beitrag stellen wir zwei Möglichkeiten der Berücksichtigung von Kovariablen im IDM gegenüber: Bei der *Discrete Event Simulation* (DES) werden relevante Ereignisse je nach Ausprägung der Kovariablen auf Individualebene simuliert, während man bei der *Effective Rate Method* (ERM) das Produkt aus Dichtefunktion der Kovariablen und Übergangsrate zu einer "effektiven Übergangsrate" ausintegriert. Die effektiven Übergangsraten können direkt in der partiellen Differentialgleichung berücksichtigt werden.

DES und ERM werden am Beispiel der Wirkung des Body Mass Index als Kovariable auf Inzidenz und Mortalität bei Typ 2 Diabetes illustriert.

[1] Brinks R, Landwehr S (2014) Theo Popul Biol 92: 62-8

Sequences of prevalence surveys with dependent data – new insights based on spatial statistics

Annika Hoyer, Ralph Brinks

Deutsches Diabetes-Zentrum (DDZ), Leibniz-Zentrum für Diabetes Forschung an der
Heinrich-Heine-Universität Düsseldorf,
Institut für Biometrie und Epidemiologie
annika.hoyer@ddz.uni-duesseldorf.de

Background

A sequence of surveys about the age-specific prevalence of a chronic disease may be represented in a Lexis diagram (Keiding 2006) – a map-like illustration of the relevant variables (calendar time and age).

In some studies, the data points in the Lexis diagram are not independent, for instance, because they contain data from people surveyed repeatedly. An example is the SHARE study, a European open cohort-study conducted in several waves. We use the SHARE study to estimate the prevalence of diabetes. Actually, we are not aware of approaches addressing the problem of modelling dependencies while estimating temporal trends of the age-specific prevalence.

Methods

To model the dependencies in the Lexis diagram, we propose a linear mixed effects (LME) model where the correlation of the random effects mimics the situation of spatial data. For comparison, we present an approach based on Kriging with thin-plate splines (TPS). Both approaches lead to estimated Lexis diagrams.

Results

In contrast to TPS, the LME yields standard errors and confidence intervals for the predicted prevalences. For the SHARE study, we found the maximal prevalence of diabetes at the age of 80 and a temporal trend where the prevalence is rising with increasing calendar time.

Conclusion

For the first time, we present methods that allow the estimation of prevalence trends incorporating dependencies in the data.

Reference

Keiding N. (2006). Event history analysis and the cross-section. *Statistics in Medicine* 25(14):2343–2364