

Prospektive Risikoanalyse: Die Ähnlichkeit von Medikamentennamen in der Drugbank-Datenbank

Prospective risk analysis: The similarity of drug names in the Drugbank database

Abstract

In various studies, a frequency of adverse drug events in medical treatments was found ranging between 0.7% and 6.5%. About 12% of these medication errors are due to confusion resulting from similar sounding names (sound-alike) or similar spellings (look-alike). The problem of sound-alike and look-alike drug names, which however should always be a unique name, is monitored by various institutions at international, European and national levels.

In order to prevent the confusion of drugs, it is recommended to maintain lists of look-alike and sound-alike drug names (LASA list). These LASA lists, however, reflect only on pairs of drug names which already had been confused

In this study, we prospectively describe the similarity of drug names based on different algorithms. For this purpose, methods analysing both orthographic and phonetic similarity were applied. The study included 62.354 drug and product names of the open accessible database „Drugbank“. According to our results about 50% of the drug names listed in the Drugbank are characterized by comparable similarity measures as the already confused drug names listed in the LASA-list of the Institute for Safe Medication Practices (ISMP) and the Federal Association of German Hospital Pharmacists e.V. (ADKA). 4.822 name pairs generated from the Drugbank were different only in one letter, while 22.039 name pairs differed in two letters or editing steps (Levenshtein distance). Analysing the phonetic similarity resulted in a resemblance between 20% and 96% of names.

Keywords: patient safety, sound alike, look-alike, prospective analysis

Zusammenfassung

In verschiedenen Untersuchungen wurde eine Häufigkeit unerwünschter Arzneimittelereignisse bei medizinischen Behandlungen zwischen 0,7% und 6,5% ermittelt. Etwa 12% dieser Fehler basieren auf Verwechslungen wegen des ähnlich klingenden Namens (sound-alike) oder des ähnlichen Aussehens. Die eindeutige Namensgebung wird auf internationaler, europäischer und nationaler Ebene von verschiedenen Institutionen überwacht. Zur Prävention dieser Verwechslungen wird empfohlen, Listen mit dokumentierten look-alike und sound-alike Verwechslungen zu führen (LASA-Liste). Die LASA-Listen geben nur die Medikamentennamenspaare wieder, wo bereits eine Verwechslung aufgetreten ist. Im Rahmen dieser Untersuchung wurde eine prospektive Studie durchgeführt, um die Ähnlichkeit von Medikamentennamen zu beschreiben. Dazu wurden sowohl Verfahren zur orthographischen als auch phonetischen Ähnlichkeit verwendet. In die Studie wurden 62.354 Wirkstoff- und Produktnamen der frei zugänglichen Datenbank Drugbank eingeschlossen.

Etwa die Hälfte der Namen weisen vergleichbare Ähnlichkeitsmaße wie die Namen der LASA-Listen des Institute for Safe Medication Practices

Thomas Schrader¹

Laura Tetzlaff¹

Cornelia Schröder¹

Eberhard Beck¹

¹ Technische Hochschule Brandenburg, Fachbereich Informatik und Medien, Brandenburg, Deutschland

(ISMP) und des Bundesverbandes Deutscher Krankenhausapotheker e.V. (ADKA) auf. 4.822 Namenspaare unterscheiden sich nur in einem, 22.039 Paare in zwei Buchstaben bzw. Editierschritten (Levenshtein-Distanz). Die phonetische Ähnlichkeit zeigt eine Ähnlichkeit zwischen 20% und 96% der Namen.

Schlüsselwörter: Patientensicherheit, sound-alike, look-alike, prospektive Analyse

Hintergrund

In einer Reihe von Untersuchungen wird die Häufigkeit von unerwünschten Arzneimittelereignissen zwischen 0,7 und 6,5% angegeben [1]. In etwa 12% der Medikationsfehler kam es zu Verwechslungen auf Grund des ähnlichen klingenden Namens (sound-alike) oder des ähnlichen Aussehens (look-alike) [2].

Über die Eindeutigkeit der Wirkstoffnamen wacht international die WHO [3]. Auf europäischer Ebene gibt es die Name Review Group der European Medicines Agency [4] und auf nationaler Ebene das Bundesamt für Arzneimittel und Medizinprodukte, die auf eine eindeutige Namensgebung achtet [5].

Zur Prävention von look-alike und sound-alike (LASA) Verwechslungen werden unterschiedliche Strategien angewandt. Es wird zum Beispiel empfohlen, Listen mit stattgefundenen Verwechslungen zu führen und die medizinischen MitarbeiterInnen darüber zu informieren. Die Federal Drug Administration (FDA) und das Institute for Safe Medication Practices (ISMP) hat die letzte Liste mit LASA-Verwechslungen 2016 veröffentlicht [6]. Die Liste der LASA-Medikamente des Bundesverbandes Deutscher Krankenhausapotheker e.V. (ADKA) geht auf das Jahr 2015 zurück [7].

Aus den Listen wird erkennbar, dass die Begriffe der sound-alike und look-alike Medikamente sich auf das Ereignis beziehen, wann die Verwechslung geschah: Trat der Fehler bei einer mündlichen Informationsübergabe auf, werden die entsprechenden Medikamente den sound-alike Namensverwechselungen zugeordnet. Verwechslungen basierend auf dem look-alike Problem können allerdings sehr unterschiedliche Gründe haben: die orthographische Ähnlichkeit der Namen – wegen der Namensähnlichkeit werden die Namen falsch gelesen – oder die morphologische Ähnlichkeit – die Ähnlichkeit des Aussehens der Handelspackung oder die Ähnlichkeit der Abpackung (z.B. Ampullen, Blister) – und schließlich die Ähnlichkeit der Darreichungsform – die Tabletten- oder Drageeform. Im Rahmen dieser Arbeit werden nur die orthographische und phonetische Ähnlichkeit untersucht, die sowohl zu sound-alike als auch look-alike Problemen führen können.

Eine Reihe von Ursachen und beeinflussenden Faktoren für diese Form der Verwechslungen wurden untersucht. Sie lassen sich individuellen, technologischen, Umwelt- und sonstigen Faktoren zuordnen [8]. Zum Beispiel spielt der Ausbildungsstand des medizinischen Personals eine wesentliche Rolle, wie häufig ähnliche Medikamente si-

cher unterschieden wurden [9]. Unter Laborbedingungen konnten bei Kognitions- und Gedächtnistest vergleichbare Fehlerraten bei Medikamentenverwechslungen wie unter realen Bedingungen nachgewiesen werden [10], was auf vergleichbare Mechanismen hinweist. Die orthographische und phonetische Ähnlichkeit wird als ein wesentlicher Faktor für die Verwechslung von Medikamentennamen angesehen [11]. So wurden auch orthographiebasierte Ähnlichkeitsmaße untersucht. Die bekanntesten Maße sind der Levenshtein-Index als Editierabstand (Anzahl der Schritte, die benötigt wird, um von einem Wort zum anderen zu gelangen) [12] und die Ähnlichkeit von Bi- und Trigrammen (Silben aus zwei bzw. drei Buchstaben) [13].

In der vorliegenden Studie wurde zum ersten Mal prospektiv eine frei zugängliche Medikamentendatenbank (Drugbank – <https://www.drugbank.ca/>) vollständig bezüglich der Häufigkeit von Ähnlichkeiten von Wirkstoff- und Produktnamen und die Eigenschaften deren Ähnlichkeit näher untersucht. Dazu wurden aus den verfügbaren LASA-Listen die Parameter zur Beschreibung der Ähnlichkeit der bereits identifizierten Namensverwechslungen ermittelt. Mit diesen Ähnlichkeitsmaßen wurde die Medikamentendatenbank untersucht, um potentielle Verwechslungskandidaten zu identifizieren. In den bisherigen Arbeiten, die auf die algorithmische Bestimmung von LASA-Medikamenten eingingen, blieben die unterschiedlichen Muster der Ähnlichkeit unberücksichtigt [12], [14], [15], [16].

Material und Methoden

Für die Analyse wurde die Datenbank „Drugbank“ verwendet, die als XML-Datei heruntergeladen werden kann [17]. Sie enthält Wirkstoffnamen, zu den Wirkstoffen die ATC-Kodierung sowie zu jedem Wirkstoff auch die internationalen Produktnamen. Stand 23.12.2017 enthielt die Datenbank 62.354 eindeutige Wirkstoff- und Produktnamen. Mehr als 1,9 Milliarden Paarvergleiche bezüglich der Ähnlichkeit wurden durchgeführt.

Es wurden Evaluationsmethoden der orthographischen und der phonetischen Ähnlichkeit verwendet. Der Algorithmus wurde in Python 3.6 unter Einschluss der Bibliotheken Distance [18], Fuzzy [19], Matplotlib [20] geschrieben. Wegen des hohen technischen Aufwands wurde der Analyseprozess für eine parallele Bearbeitung auf mehreren Rechenkernen optimiert [21].

Methoden der orthographischen Ähnlichkeit

Folgende orthographischen Ähnlichkeitsmaße wurden für diese Evaluation berechnet:

- Die Levenshtein-Distanz ist die einfache Editier-Distanz und zählt, wie viele Buchstaben durch Vertauschen, Einfügen und Löschen geändert werden müssen, um von einem Wort zum anderen zu gelangen [22]. Zum Beispiel: Um vom Wort „Husten“ zum Wort „Hasten“ zu gelangen, muss lediglich der Buchstabe „u“ durch „a“ ersetzt werden – Die Levenshtein-Distanz beträgt 1. Die Wörter sind sehr ähnlich.
- Die normierte Levenshtein-Distanz (Levenshtein-Index) basiert auf der Normierung auf die Länge des kürzesten Wortes [18]. Die Levenshtein-Distanz wird durch die Anzahl der Buchstaben des kürzeren Wortes dividiert. In diesem Fall sind beide Wörter gleich lang: $1/6=0,16667$. Die Ähnlichkeit ist auch hier groß. Je kleiner der Levenshtein-Index ist, desto ähnlicher sind sich die Wörter. Der Levenshtein-Index würde noch kleiner werden, je länger die Worte sind. Damit berücksichtigt dieser Index die Wortlänge. Kleine Levenshtein-Distanzen in sehr langen Wörtern führen zu einer hohen Ähnlichkeit im Index-Wert.
- Die Jaccard-Distanz berechnet das Verhältnis aus der gemeinsamen Schnittmenge und der Vereinigung der zu vergleichenden Textstrukturen [23]. Schematisch sieht das beim „Husten/Hasten“-Beispiel so aus (die Anzahl der gleichen Buchstaben kommt in den Zähler, die Anzahl aller verwendeten Buchstaben in den Nenner):

$$1 - \frac{H_{s,t,e,n}}{H_{u,a,s,t,e,n}} = 1 - \frac{5}{7} = 1 - 0.7142 = 0.2858$$

Es wurde weiterhin die längste gemeinsame Buchstabenkette (LCS – longest common substring) [24] bestimmt.

Methoden der phonetischen Ähnlichkeit

Für die Untersuchung der phonetischen Ähnlichkeit ist das Soundex-Verfahren das bekannteste, da es in den Datenbank-Suchalgorithmen implementiert wurde [25]. Für diese Untersuchung wurden zwei Weiterentwicklungen verwendet:

1. Der sog. Double Metaphone Algorithmus kodiert Worte in phonetische Einheiten konstanter Länge um und erlaubt darüber den Vergleich [26].
2. Der New York State Identification and Intelligence System-Algorithmus (NYSIIS) basiert auf einer phonetischen Kodierung des jeweiligen Wortes [27].

Ergebnisse

Etwa 80% der Namenszwillinge auf den LASA-Listen wiesen eine Wortlängendifferenz von <4 und einen Levenshtein-Distanz von <6 auf. Diese wurden als Schwell-

lenwerte für die Analyse der Drugbank herangezogen, die 62.354 Namen für Wirkstoffe und Produkte enthält. 1,9 Milliarden Paarvergleiche wurden durchgeführt. 251.040 Paare mit 29.887 Namen erfüllten die genannten Bedingungen.

Bezüglich der Wortlängen unterscheiden sich Wirkstoffnamen und Produktnamen deutlich (Abbildung 1): Wirkstoffnamen haben eine Wortlänge von durchschnittlich 11,9 Buchstaben (Standardabweichung 4,2, Minimum 4 Buchstaben, Maximum 60 Buchstaben). Produktnamen verwenden im Namen durchschnittlich 26,1 Buchstaben (Standardabweichung 18,2, Minimum 2, Maximum 367 Buchstaben).

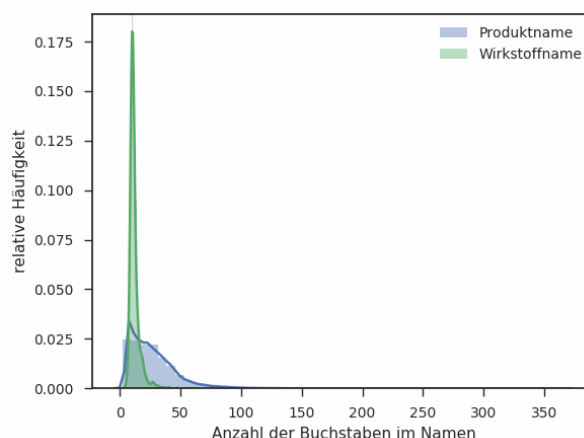


Abbildung 1: Verteilung der relativen Häufigkeit der Wortlänge (Anzahl der Buchstaben) der Namen

Die orthographische Ähnlichkeit

Bei den 1,9 Milliarden Paarvergleichen wurden alle Paare ausgeschlossen, deren Levenshtein-Distanz >5 war. Die verbleibenden 251.040 Paare haben am häufigsten einen Editierabstand von drei (92.411 Paare). 4.822 Paare unterscheiden sich nur in einem, 22.039 Paare in zwei Buchstaben bzw. Editierschritten (Abbildung 2).

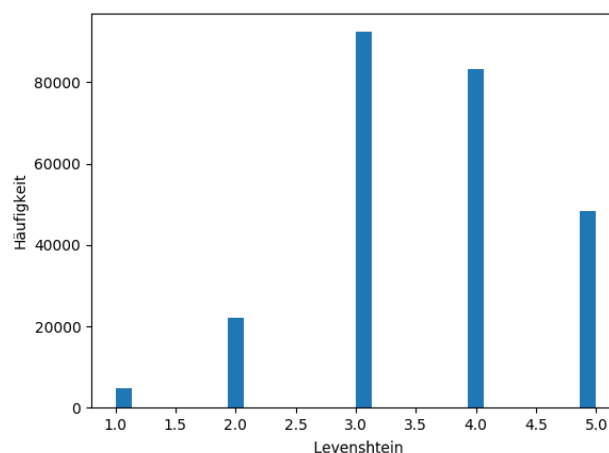


Abbildung 2: Verteilung der Levenshtein-Distanz in 251.040 Namenspaaren

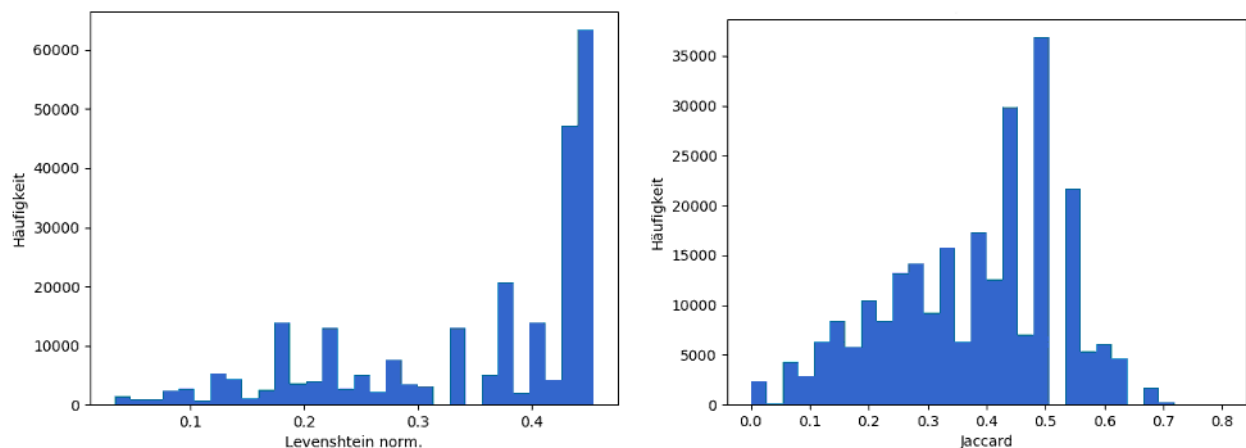


Abbildung 3: Häufigkeitsverteilungen des normierten Levenshtein- und des Jaccard-Indexes

Zwei typische Beispiele für Paare mit einem Levenshtein-Maß von 1:

- atropine injection bp 0.6mg/ml – atropine injection bp 0.4mg/ml
- bicarbonate concentrate d16000 – bicarbonate concentrate d17000

Es zeigt sich, dass hier besonders häufig die geringen Unterschiede in den Dosisangaben liegen, die hier Bestandteil des Namens waren (zahlreiche andere Namen beinhalten keine Dosisangaben).

Bei einer Editierdistanz von zwei spielt die reine Ähnlichkeit im Namen eine größere Rolle. Die folgende Beispielliste zeigt auch, dass ein Medikamentenname ähnlich auch zu mehreren anderen Namen sein kann:

- niamid – niazin
- niamid – tisamid
- niamid – nialamid
- niamid – niamidal

Die Verteilung des normierten Levenshtein folgt keiner Normalverteilung. Es erfolgt ein sprunghafter Anstieg der Häufigkeit von Werten $>0,43$. Der Levenshtein-Index korreliert nicht mit der Levenshtein-Distanz. Der Jaccard-Index hat eine Spannweite von 0 bis 0,73 und korreliert ebenfalls nicht mit der Levenshtein-Distanz (Abbildung 3). Das längste gemeinsame Segment (longest common sequence – LCS) wurde für jedes Namenspaar bestimmt. In der weiteren Analyse wurden alle LCS ausgeschlossen, die eine Länge von 1 Buchstaben aufwiesen. Es wurden also alle n-Grame betrachtet mit $n > 2$ (Bigram, Trigram ect.). Es kann vorkommen, dass mehrere LCS in einem Namenspaar zu finden sind, weshalb die Gesamtanzahl der LCS die Anzahl der untersuchten Paare überschreiten kann.

Die Länge der n-Grame ist sehr unterschiedlich, am häufigsten treten Bi- und Trigramme auf, es gibt ein lokales Maximum von bei einer Länge von 11 und 17 Buchstaben (Abbildung 4). Dabei ist zu berücksichtigen, dass solche Längen der n-Grame vor allem dann zustande kommen, wenn das Levenshtein-Maß klein ist (<3) und sich die Unterschiede überwiegend in den Dosisangaben abspielen. Am häufigsten kommen die LCS ‚ine‘ (9.012 Mal),

‚in‘ (11.011 Mal) und ‚acid concentrate‘ (18.350 Mal) vor.

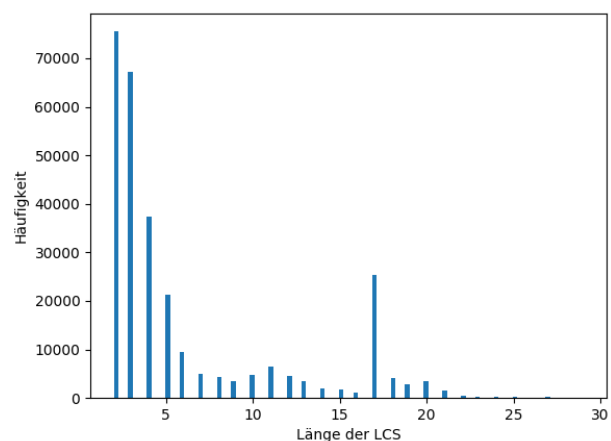


Abbildung 4: Häufigkeitsverteilung der Länge der gemeinsamen Segmente

Die phonetische Ähnlichkeit

Die drei verwendeten phonetischen Verfahren kommen zu sehr unterschiedlichen Ergebnissen. Hier wurden alle 62.354 Wirkstoff- und Produktnamen in die Analyse einbezogen.

Der DMetaphone-Algorithmus versucht für jedes Wort zwei Phonem-Hashwerte auszurechnen, wobei der erste Wert (DMetaphone 1) weniger streng bezüglich der Phonembildung ist als der zweite (DMetaphone 2). Nicht für alle Namen kann ein Phonem-Hash-Wert ermittelt werden.

Für DMetaphone 1 können für acht Namen keine Hashwerte berechnet werden. Nur 2.497 Wirkstoff- oder Produktnamen (4% aller Namen) sind nach diesem Algorithmus eindeutig. Die ermittelten Hashwerte sind durchschnittlich zu 11 Namen gleich. Ein Phonem-Hashwert tritt 538 Mal auf. Beispiele für gleiche DMetaphone 1-Hashwerte sind:

- ranidil: RNTL
- ran-tolterodine: RNTL
- reneidil srt 10mg: RNTL

- ran-duloxetine: RNTL
- rentylin: RNTL
- renedil: RNTL

Phonem-Hashwerte des DMetaphone 2-Algorithmus können nur für 6.864 Namen berechnet werden, davon sind 5.997 Hashwerte nicht eindeutig (87,3% aller mit DMetaphone 2 analysierten Namen). Durchschnittlich sind die Phonem-Hashwerte zu 7 Namen gleich. Einmal sind 230 Namen phonetisch gleich. Beispiele für gleiche DMetaphone 2-Hashwerte sind:

- erythro-ec: ARTR
- arthritic pain relief: ARTR
- erythropoietin: ARTR
- erythromycin: ARTR
- erythrocin liq 250: ARTR
- arthrexin: ARTR

Der NYSIIS-Algorithmus erlaubt die Berechnung von Hashwerten für alle Namen. 12.221 Namen sind phonetisch nicht eindeutig (19,5% aller Wirkstoff- und Produkt-namen). Durchschnittlich treten die Phonem-Hashwerte 2,6 Mal auf, ein Hashwert ist für 137 Namen gleich. Beispiele für gleiche NYSIIS-Hashwerte sind:

- amikin: ANACAN
- anacin es: ANACAN
- anacaine: ANACAN
- anacin-3: ANACAN
- amukin: ANACAN
- anacin: ANACAN
- amikin: ANACAN

Diskussion

Bisher wurde keine prospektive, vollständige Untersuchung auf phonetische und orthographische Ähnlichkeit einer internationalen Medikamentendatenbank durchgeführt. Die hier verwendeten Schwellenwerte basieren auf der Analyse der Werte aus den LASA-Listen und spiegeln wider, dass die gefundenen Namensähnlichkeiten vergleichbare Eigenschaften haben wie die in den LASA-Listen aufgeführten Namen.

In der Studie von Rash-Foanio et al. wurde im klinischen Betrieb versucht, basierend auf einer Orthographie-Ähnlichkeit sound-alike und look-alike Fehler in den Verordnungen unter Einbeziehung der Diagnose zu identifizieren [15]. Allerdings greift diese Arbeit die grundsätzlichen Unterschiede von phonetischer, orthographischer und morphologischer Ähnlichkeit nicht auf. Bryan et al. haben 2015 die formalen und semantischen Eigenschaften der WHO-Liste der internationalen nicht-proprietären Wirkstoffnamen durchgeführt.

Die Untersuchung der internationalen Datenbank Drugbank zeigt, dass bei fast der Hälfte der Wirkstoff- und Produkt-namen (national und international) ein erhöhtes Verwechslungspotential besteht. Das drückt sich in einem Levenshtein-Index von <6 und einer Wortlängendifferenz von ≤4 aus. Die Ähnlichkeit der Namen kann mit den

anderen orthographischen und phonetischen Ähnlichkeitsmaßen noch unterstrichen werden. Besonders problematisch dabei ist, dass eine Ähnlichkeit nicht nur zu einem anderen Wirkstoff- oder Produkt-namen sondern bis zu max. 236 anderen Namen (Median 7, Modalwert 1) besteht.

Insbesondere bei den kurzen Editierdistanzen von 1 oder 2, wo sich die Unterschiede vor allem in den Dosis-Angaben widerspiegeln, kommt dem Aussehen der Verpackung eine große Bedeutung zu.

Wenn Herstellernamen Bestandteil des Produkt-namens werden, hat das aus analytischer Sicht zwei Konsequenzen: Die Untersuchung der gemeinsamen Wortanteile (LCS) zeigt lange gemeinsame Wortsegmente, die dann auch rechnerisch eine hohe orthographische Ähnlichkeit ergeben. Unter diesen Bedingungen sind dann nach LCS von mehr als 10 Buchstaben möglich.

Abhängig von der verwendeten Methode der Analyse der phonetischen Ähnlichkeit, sind zwischen 20% (NYSIIS-Methode) und 96% (DMetaphone 1-Methode) nicht eindeutig. Bedingt durch den Algorithmus von DMetaphone 1, der unabhängig der Namenslänge alles auf eine vier Buchstaben lange phonetische Repräsentation reduziert, ist dessen Trennschärfe schlechter: es werden mehr Namen als ähnlich identifiziert. Damit sind die Ergebnisse weniger spezifisch.

Diese Studie zeigt, dass viele LASA-Probleme hausgemacht sind und sich vermeiden lassen, wenn bei der Vergabe der Namen strengere Regeln auf zumindest orthographische Ähnlichkeit umgesetzt werden würden. Schon Lambert et al. fordern 2005 eine systematische Untersuchung der Ähnlichkeit vor der Zulassung von Medikamenten [28]. Die vorliegende Studie zeigt zum Beispiel auf, welche Bedeutung die Dosisangaben im Namen haben, die zu einer Verwechslung führen können. Die Trennung der Dosisangabe und die optisch getrennte Darstellung nach einem (noch zu entwickelndem) Muster könnte die Verwechslungsgefahr verringern. Die Regierungsorganisation Health Canada nutzt einen sog. Brand Name Assessment and Review Process bei dem auch ein phonetischer und orthographischer Ähnlichkeitsscore berechnet wird, der <50% sein sollte [29]. Allerdings wird auf die Art der Berechnung nicht näher eingegangen.

Prospektiv lassen sich diese Probleme mit den aktuellen Analysemethoden berechnen und es kann versucht werden, bei der Zulassung sowie in der Zeit danach Änderungen des Namens zu bewirken. Während es bei Geräteherstellern möglich ist, dass sie Teile so zu verändern haben, dass keine Verwechslungen auftreten können [30], muss bei den Medikamenten noch ein Umdenken erfolgen.

In diesem Zusammenhang scheint es sinnvoll zu sein, die bisherige Einteilung in sound-alike und look-alike Kategorien soweit zu ergänzen, dass weiter differenziert werden kann, ob es sich um eine orthographische, phonetische oder morphologische Ähnlichkeit handelt. Es wurden keine Studien gefunden, die zwischen den unterschiedlichen Formen der Ähnlichkeit differenzieren und dafür getrennt Häufigkeiten der Verwechslungen angeben.

Die orthographische und phonetische Ähnlichkeit kann gut untersucht und beschrieben werden, da dafür ein großes Methodenrepertoire zur Verfügung steht [22], [24]. Ein schrittweises Vorgehen bei der Analyse scheint dabei sinnvoll zu sein: die Editierdistanz (Levenshtein-Distanz) ist zunächst ein starkes Indiz für die Ähnlichkeit. Allerdings kann die Ähnlichkeit nur im Zusammenhang mit der Wichtung zur Länge mittels Levenshtein-Index und Jaccard-Index eingehender beurteilt werden. In weiteren Untersuchungen sollte analysiert werden, welche Bedeutung die Position des gemeinsamen Wortanteils (LCS) hat. Wenn der gemeinsame Wortteil auch noch an gleicher Position vorkommt, dann ist davon auszugehen, dass die Ähnlichkeit höher zu bewerten ist. Schwieriger dagegen ist die Beurteilung der morphologischen Ähnlichkeit: sie umfasst unterschiedliche Ebenen (Handelsverpackung, Abpackung und Darreichungsform) und bedarf bildanalytischer Methoden. Bisher ist die morphologische Ähnlichkeit nicht systematisch untersucht worden. Obwohl die Anzahl der ähnlichen Medikamente sehr hoch ist, muss einschränkend betont werden, dass die Datenbank Drugbank auch zahlreiche eher den kosmetischen und pflegerischen Produkten zuzuordnenden Substanzen enthält. Die vorliegende Untersuchung kann das Auftreten von Verwechslungen nicht exakt vorhersagen, da dies von verschiedenen Faktoren abhängt. Die berechneten Ähnlichkeitsmaße stehen derzeit noch für sich allein, ohne in den Kontext der Darreichungsform und des Indikationsgebietes eingebunden zu sein. In weiteren Studien soll ein Instrumentarium entwickelt werden, das die Ähnlichkeit der Medikamente besser charakterisiert und hinsichtlich verschiedener Kriterien wichtet. Für die Untersuchung der morphologischen Ähnlichkeit müssen Datenbanken angelegt werden, die das Aussehen der Verpackung und Darreichungsform speichern und damit einer systematischen Analyse zugänglich machen.

Anmerkung

Interessenkonflikte

Die Autoren erklären, dass sie keine Interessenkonflikte in Zusammenhang mit diesem Artikel haben.

Literatur

1. von Laue NC, Schwappach DL, Koeck CM. The epidemiology of preventable adverse drug events: a review of the literature. *Wien Klin Wochenschr.* 2003 Jul;115(12):407-15. DOI: 10.1007/BF03040432
2. Leape LL, Bates DW, Cullen DJ, Cooper J, Demonaco HJ, Gallivan T, Hallisey R, Ives J, Laird N, Laffel G, et al. Systems analysis of adverse drug events. ADE Prevention Study Group. *JAMA.* 1995 Jul 5;274(1):35-43. DOI: 10.1001/jama.1995.03530010049034
3. World Health Organisation. Lists of Recommended and Proposed INNs. [accessed 2017 Dec 23]. Available from: <http://www.who.int/medicines/publications/druginformation/innlists/en/>
4. European Medicines Agency. (Invented) Name Review Group. [accessed 2017 Dec 23]. Available from: http://www.ema.europa.eu/ema/index.jsp?curl=pages/contacts/CHMP/people_listing_000035.jsp&murl=menus/about_us/about_us.jsp&mid=WC0b01ac0580028dd4&jsearch=true
5. Hahnenkamp C, Rohe J, Thomeczek C. Patientensicherheit: Ich sehe was, was du nicht schreibst ... *Dtsch Arztebl.* 2011;108(36):A 1850-4.
6. Institute for Safe Medication Practices (ISMP). Look-Alike Drug Names with Recommended Tall Man Letters. 2016. Available from: <https://www.ismp.org/recommendations/tall-man-letters-list>
7. Bundesverband Deutscher Krankenhausapotheker e.V. (ADKA). Dokumentation Pharmazeutischer Interventionen im Krankenhaus (ADKA-DokuPIK) – Tabellen SA/LA. 2015 [accessed 2018 Feb 03]. Available from: https://www.adka-dokupik.de/index.cfm?CFID=5816305&CFTOKEN=18720097&pt=Down_Cat_Liste&cat_id=72A855FD-E227-E85F-6960345B971E7686
8. Kawano A, Li Q, Ho C. Preventable Medication Errors – Look-alike/Sound-alike Drug Names. *Pharm Connect.* 2014;28-33.
9. Tsuji T, Irisa T, Tagawa S, Kawashiri T, Ikeshue H, Kokubu C, Kanaya A, Egashira N, Masuda S. Differences in recognition of similar medication names between pharmacists and nurses: a retrospective study. *J Pharm Health Care Sci.* 2015 Jul 7;1:19. DOI: 10.1186/s40780-015-0017-4
10. Schroeder SR, Salomon MM, Galanter WL, Schiff GD, Vaida AJ, Gaunt MJ, Bryson ML, Rash C, Falck S, Lambert BL. Cognitive tests predict real-world errors: the relationship between drug name confusion rates in laboratory-based memory and perception tests and corresponding error rates in large pharmacy chains. *BMJ Qual Saf.* 2017 May;26(5):395-407. DOI: 10.1136/bmjqs-2015-005099
11. Lambert BL, Chang KY, Lin SJ. Effect of orthographic and phonological similarity on false recognition of drug names. *Soc Sci Med.* 2001 Jun;52(12):1843-57. DOI: 10.1016/S0277-9536(00)00301-4
12. Bryan R, Aronson JK, ten Hacken P, Williams A, Jordan S. Patient Safety in Medication Nomenclature: Orthographic and Semantic Properties of International Nonproprietary Names. *PLoS One.* 2015 Dec 23;10(12):e0145431. DOI: 10.1371/journal.pone.0145431
13. Lambert BL. Predicting look-alike and sound-alike medication errors. *Am J Health Syst Pharm.* 1997 May 15;54(10):1161-71.
14. Kovacic L, Chambers C. Look-alike, sound-alike drugs in oncology. *J Oncol Pharm Pract.* 2011 Jun;17(2):104-18. DOI: 10.1177/1078155209354135
15. Rash-Foano C, Galanter W, Bryson M, Falck S, Liu KL, Schiff GD, Vaida A, Lambert BL. Automated detection of look-alike/sound-alike medication errors. *Am J Health Syst Pharm.* 2017 Apr 1;74(7):521-7. DOI: 10.2146/ajhp150690

16. Kondrak G, Dorr B. Identification of Confusable Drug Names: A New Approach and Evaluation Methodology. In: COLING '04: Proceedings of the 20th international conference on Computational Linguistics; 2004 Aug 23-27; Geneva, Switzerland. Stroudsburg, PA, USA: Association for Computational Linguistics; 2004. Article No. 952. DOI: 10.3115/1220355.1220492
17. Wishart DS, Feunang YD, Guo AC, Lo EJ, Marcu A, Grant JR, Sajed T, Johnson D, Li C, Sayeeda Z, Assempour N, Iynkkaran I, Liu Y, Maciejewski A, Gale N, Wilson A, Chin L, Cummings R, Le D, Pon A, Knox C, Wilson M. DrugBank 5.0: a major update to the DrugBank database for 2018. *Nucleic Acids Res.* 2018 Jan 4;46(D1):D1074-D1082. DOI: 10.1093/nar/gkx1037
18. Meyer M. Distance: Utilities for comparing sequences. Python Package Index (PyPI). Available from: <https://pypi.org/project/Distance/>
19. Fuzzy: Fast Python phonetic algorithms. Python Package Index (PyPI). Available from: <https://pypi.org/project/Fuzzy/>
20. Matplotlib: Python plotting – Matplotlib 2.1.1 documentation. [accessed 2018 Feb 03]. Available from: <https://matplotlib.org/>
21. Multiprocessing – Process-based parallelism – Python 3.6.6 documentation. [accessed 2018 Feb 03]. Available from: <https://docs.python.org/3/library/multiprocessing.html>
22. Navarro G. A Guided Tour to Approximate String Matching. *ACM Comput Surv.* 2001;33:31-88. DOI: 10.1145/375360.375365
23. Deng F, Siersdorfer S, Zerr S. Efficient Jaccard-based Diversity Analysis of Large Document Collections. In: Proceedings of the 21st ACM international conference on Information and knowledge management. New York: ACM; 2012. p. 1402-11.
24. Bergroth L, Hakonen H, Raita T. A survey of longest common subsequence algorithms. In: Proceedings of the Seventh International Symposium on String Processing and Information Retrieval – SPIRE 2000; 2000 Sep 27-29; A Coruna, Spain. IEEE; 2000. p. 39-48. DOI: 10.1109/SPIRE.2000.878178
25. Holmes D, McCabe MC. Improving precision and recall for Soundex retrieval. In: Proceedings of the International Conference on Information Technology: Coding and Computing; 2002 Apr 8-10; Las Vegas, NV, USA. IEEE; 2002. p. 22-6. DOI: 10.1109/ITCC.2002.1000354
26. Phillips L. The Double Metaphone Search Algorithm. Dr. Dobb's; 2000 [accessed 2018 Feb 03]. Available from: <http://www.drdoobbs.com/the-double-metaphone-search-algorithm/184401251>
27. NYSIIS. [accessed 2018 Feb 03]. Available from: <https://xlinux.nist.gov/dads/HTML/nysiis.html>
28. Lambert BL, Lin SJ, Tan H. Designing safe drug names. *Drug Saf.* 2005;28(6):495-512. DOI: 10.2165/00002018-200528060-00003
29. Health Canada, editor. Guidance Document for Industry – Review of Drug Brand Names. Ottawa: Minister of Health; 2014 [accessed 2017 Dec 23]. Available from: <https://www.canada.ca/en/health-canada/services/drugs-health-products/reports-publications/medeffect-canada/guidance-document-industry-review-drug-brand-names.html>
30. Luer-Konnektoren: Verwechslungen vermeiden. Apotheke + Marketing. 2017 Jan 31 [accessed 2018 Feb 20]. Available from: <https://www.apotheke-und-marketing.de/luer-konnektoren-verwechslungen-vermeiden-1873116.html>

Erratum

Im ersten und zweiten Satz von Abstract, Zusammenfassung und Hintergrund wurden die Prozentangaben korrigiert. Unter Hintergrund wurden zudem die zugehörigen Quellenangaben angepasst. Durch das Ergänzen einer Quelle hat sich die Zählung der nachfolgenden Literaturverweise verschoben und das Literaturverzeichnis geändert.

Korrespondenzadresse:

Prof. Dr. Thomas Schrader
Technische Hochschule Brandenburg, Fachbereich Informatik und Medien, 14770 Brandenburg, Deutschland
schrader@th-brandenburg.de

Bitte zitieren als

Schrader T, Tetzlaff L, Schröder C, Beck E. Prospektive Risikoanalyse: Die Ähnlichkeit von Medikamentennamen in der Drugbank-Datenbank. *GMS Med Inform Biom Epidemiol.* 2018;14(2):Doc09.
DOI: 10.3205/mibe000187, URN: urn:nbn:de:0183-mibe0001873

Artikel online frei zugänglich unter

<http://www.egms.de/en/journals/mibe/2018-14/mibe000187.shtml>

Veröffentlicht: 30.08.2018

Veröffentlicht mit Erratum: 18.04.2019

Copyright

©2018 Schrader et al. Dieser Artikel ist ein Open-Access-Artikel und steht unter den Lizenzbedingungen der Creative Commons Attribution 4.0 License (Namensnennung). Lizenz-Angaben siehe <http://creativecommons.org/licenses/by/4.0/>.