



universität
wien

MASTERARBEIT / MASTER'S THESIS

Titel der Masterarbeit / Title of the Master's Thesis

„Arzt-Patient-Kommunikation mit SayHi Translate

Eine aufgabenbezogene Evaluierung im Sprachenpaar Deutsch-Ungarisch“

verfasst von / submitted by

Julia Glocknitzer, BA

angestrebter akademischer Grad / in partial fulfilment of the requirements for the degree of
Master of Arts (MA)

Wien, 2020 / Vienna 2020

Studienkennzahl lt. Studienblatt /
degree programme code as it appears on
the student record sheet:

A 070 381 331

Studienrichtung lt. Studienblatt /
degree programme as it appears on
the student record sheet:

Masterstudium Translation Ungarisch Deutsch

Betreut von / Supervisor:

Univ.-Prof. Mag. Dr. Franz Pöchhacker

Danksagung

Ich möchte mich an dieser Stelle bei all denjenigen bedanken, die mich während meines Studiums und bei der Erstellung dieser Masterarbeit unterstützt haben.

Zuerst gebührt mein Dank dem Betreuer dieser Masterarbeit, Herrn Univ.-Prof. Mag. Dr. Franz Pöchlhammer, der vom ersten Moment an ein offenes Ohr für meine Ideen und Fragen hatte, mir stets mit viel Geduld und mit seinem hohen Fachwissen zur Seite stand und mich mit konstruktiver Kritik immer wieder auf den richtigen Weg führte.

Ein besonderer Dank geht an die Teilnehmenden meines Experimentes, ohne die diese Arbeit nicht hätte entstehen können. Außerdem möchte ich mich herzlich bei Frau Mag. Dr. Erna-Maria Trubel bedanken. Sie half mir nicht nur bei der Durchführung des Experimentes, sondern unterstützte und motivierte mich während meines gesamten Studiums.

Ein großes Dankeschön geht ebenfalls an Andreas Lehner, der sich Zeit genommen hat, meine Masterarbeit zu lesen, zu kommentieren und zu verbessern. Außerdem möchte ich Verena Keilen für das Korrekturlesen und Lektorat meiner Arbeit danken.

Meinen Freunden und meinem Freund danke ich besonders für den starken emotionalen Rückhalt über die Dauer meines gesamten Studiums.

Zu guter Letzt möchte ich mich bei meinen Eltern bedanken, die mich während meiner gesamten Studienzeit motiviert haben, nie aufzugeben und immer an mich selbst zu glauben. Vielen Dank für eure Liebe und Unterstützung.

Inhaltsverzeichnis

Tabellenverzeichnis	7
Abbildungsverzeichnis	9
Einleitung	10
1. Begriffserklärung.....	12
2. Geschichte des maschinellen Dolmetschens	14
2.1 Anfänge der maschinellen Übersetzung.....	14
2.2 Der ALPAC-Bericht	15
2.3 Das Wiedererwachen der Forschung	15
2.4 Die Entwicklung maschineller Dolmetsch-Systeme.....	16
3. Funktionsweise maschineller Dolmetsch-Systeme.....	20
3.1 Automatische Spracherkennung.....	20
3.1.1 Herausforderungen der automatischen Spracherkennung bei spontan gesprochener Sprache.....	22
3.2 Maschinelle Übersetzung.....	25
3.2.1 Regelbasierte Ansätze	26
3.2.2 Statistikbasierte Ansätze	29
3.2.3 Neuronale MÜ	35
3.2.4 Probleme der MÜ	37
3.3 Sprachsynthese.....	39
4. SayHi Translate	41
4.1 Überblick.....	41
4.2 User Interface	43
5. Dolmetschen im Gesundheitswesen	47
5.1 Die Rolle der DolmetscherInnen in medizinischen Settings	47
5.2 Dolmetschen im österreichischen Gesundheitswesen.....	49
6. Forschungsstand.....	53
7. Experiment.....	58
7.1 Forschungsfragen	58
7.2 Teilnehmer des Experiments	58
7.3 Ablauf der Untersuchung.....	59

7.4 Evaluierungsmethode.....	60
7.4.1 Problematik der Evaluierung von MD-Systemen.....	60
7.4.2 Aufgabenbezogene Evaluierung.....	61
7.5 Transkription des Experimentes	64
8. Ergebnisse	67
8.1 BenutzerInnenfreundlichkeit	67
8.2 Spracherkennungsfehler	70
8.2.1 Geringfügige und mittelschwere Spracherkennungsfehler	71
8.2.2 Schwerwiegende Spracherkennungsfehler.....	81
8.3 Übersetzungsfehler	88
8.3.1 Geringfügige Übersetzungsfehler	88
8.3.2 Fehler, die die Bedeutung des Satzes nicht wesentlich beeinträchtigen.....	90
8.3.3 Fehler, die das Satzverständnis merklich beeinträchtigen	93
8.3.4 Schwerwiegende Übersetzungsfehler.....	98
8.4 Usability.....	103
9. Zusammenfassung und Schlussfolgerungen	112
Bibliographie.....	116
Anhang	123
Abstract	138

Tabellenverzeichnis

Tabelle 1: Entwicklungsphasen maschineller Dolmetsch-Systeme (Waibel 2015: 107)	17
Tabelle 2: Berechnung der Übersetzungswahrscheinlichkeiten für jedes Wort im Satz „Das Haus ist klein“ nach Koehn (2010: 84)	31
Tabelle 3: Liste der medizinischen Sätze, die mit Google Translate übersetzt wurden (Patil & Davies 2014)	56
Tabelle 4: Ausschnitt aus der Transkription des Experimentes (Zeilen 1 – 5)	64
Tabelle 5: Ausschnitt aus der Transkription des Experimentes, Beispiel für einen Spracherkennungsfehler (Zeilen 46 – 49)	66
Tabelle 6: Transkription des Experimentes, Spracherkennungsfehler (Zeilen 12, 41, 54 und 68)	71
Tabelle 7: Transkription des Experimentes, Spracherkennungsfehler (Zeilen 9 und 52)	73
Tabelle 8: Transkription des Experimentes, Spracherkennungsfehler (Zeilen 27 und 30)	74
Tabelle 9: Transkription des Experimentes, Spracherkennungsfehler (Zeile 54)	75
Tabelle 10: Transkription des Experimentes, Spracherkennungsfehler (Zeilen 13, 17 und 70)	76
Tabelle 11: Transkription des Experimentes, Spracherkennungsfehler (Zeile 60)	77
Tabelle 12: Transkription des Experimentes, Spracherkennungsfehler (Zeilen 43 und 44)	77
Tabelle 13: Transkription des Experimentes, Spracherkennungsfehler (Zeilen 15 und 18)	78
Tabelle 14: Transkription des Experimentes, Spracherkennungsfehler (Zeilen 11, 18, 32 und 42)	79
Tabelle 15: Transkription des Experimentes, Spracherkennungsfehler (Zeilen 37, 38, 41 und 59)	81
Tabelle 16: Transkription des Experimentes, Spracherkennungsfehler (Zeile 15)	83
Tabelle 17: Transkription des Experimentes, Spracherkennungsfehler (Zeilen 44 und 45)	83
Tabelle 18: Transkription des Experimentes, Spracherkennungsfehler (Zeilen 46 – 49)	84
Tabelle 19: Transkription des Experimentes, Spracherkennungsfehler (Zeilen 19 – 26)	85
Tabelle 20: Transkription des Experimentes, Spracherkennungsfehler (Zeilen 60 – 63)	86
Tabelle 21: Transkription des Experimentes, Spracherkennungsfehler (Zeilen 66 – 67)	87
Tabelle 22: Beispiel für ein Gewichtungsschema nach Cap (2003: 61)	88
Tabelle 23: Transkription des Experimentes, Übersetzungsfehler, Orthographiefehler (Zeile 64)	89
Tabelle 24: Transkription des Experimentes, Übersetzungsfehler, Stilistische Fehler (Zeilen 12 und 16)	89
Tabelle 25: Transkription des Experimentes, Übersetzungsfehler, Stilistische Fehler (Zeile 61)	90
Tabelle 26: Transkription des Experimentes, Übersetzungsfehler, Falsche Präpositionen (Zeilen 30 und 34)	90
Tabelle 27: Transkription des Experimentes, Übersetzungsfehler, Deklinationsfehler (Zeilen 7 und 30)	91
Tabelle 28: Transkription des Experimentes, Übersetzungsfehler, Falsche Wortstellung (Zeilen 4 und 39)	92
Tabelle 29: Transkription des Experimentes, Übersetzungsfehler, Pronominalbezugsfehler (Zeilen 5 und 13)	93
Tabelle 30: Transkription des Experimentes, Übersetzungsfehler, Pronominalbezugsfehler (Zeilen 15 und 19)	94
Tabelle 31: Transkription des Experimentes, Übersetzungsfehler, Konjugationsfehler (Zeilen 18, 34, 35 und 68)	95
Tabelle 32: Transkription des Experimentes, Übersetzungsfehler, Englische Ausdrücke (Zeilen 10 und 67)	97
Tabelle 33: Transkription des Experimentes, Schwerwiegende Übersetzungsfehler (Zeilen 18, 44 und 65)	98
Tabelle 34: Transkription des Experimentes, Schwerwiegende Übersetzungsfehler (Zeile 42)	99
Tabelle 35: Transkription des Experimentes, Schwerwiegende Übersetzungsfehler (Zeile 32)	99

Tabelle 36: Transkription des Experimentes, Schwerwiegende Übersetzungsfehler (Zeile 58).....	100
Tabelle 37: Transkription des Experimentes, Schwerwiegende Übersetzungsfehler (Zeilen 15, 36, 51 und 64)	101
Tabelle 38: Transkription des Experimentes, nonverbale Kommunikation des Arztes während der Untersuchung	104
Tabelle 39: Transkription des Experimentes, nonverbale Kommunikation (Zeilen 1 – 4, 52 – 53)	105

Abbildungsverzeichnis

Abbildung 1: Veranschaulichung der Verschiedenheit von Text und Sprachsignal (Pfister & Kaufmann 2017: 7).....	21
Abbildung 2: Ein typischer Spracherkenner (Speech-to-Text) (Waibel 2015: 105)	22
Abbildung 3: Einflussfaktoren auf automatische Spracherkennungssysteme nach Pfister & Kaufmann (2017: 25).....	23
Abbildung 4: MÜ Vauquois’Dreieck der MÜ nach Hutchins & Somers (1992: 107).....	26
Abbildung 5: Schema für ein direktes MÜ-System (Carstensen et al. 2004: 566).....	27
Abbildung 6: Schema für Transferbasierte Übersetzung (Carstensen et al. 2010: 646).....	28
Abbildung 7: Schema für Interlingua-basierte Übersetzung (Carstensen et al. 2010: 646)	28
Abbildung 8: Schema für die statistische MÜ (Carstensen et al. 2010: 650).....	32
Abbildung 9: Neuronales Netz mit verdeckten Schichten (hidden layers) (Laki 2018: 170, übersetzt aus dem Ungarischen).....	36
Abbildung 10: Schematische Abbildung einer Schicht des rekurrenten neuronalen Netzes (Laki 2018: 171, übersetzt aus dem Ungarischen)	37
Abbildung 11: Ablauf der Sprachsynthese nach Pfister & Kaufmann (2017: 197)	40
Abbildung 12: Logo und App-Symbol von SayHi Translate (Screenshot).....	41
Abbildung 13: Sprachen und Dialekte in SayHi Translate (Screenshot).....	42
Abbildung 14: Startbildschirm von SayHi Translate (Screenshot)	43
Abbildung 15: Anweisungen für die verschiedenen Funktionen von SayHi Translate (Screenshot)	44
Abbildung 16: Einstellungen in SayHi Translate (Screenshot)	44
Abbildung 17: MD-Prozess (Screenshot).....	45
Abbildung 18: Manuelle Korrekturfunktion in SayHi Translate (Screenshot).....	46
Abbildung 19: Gesamtbewertung der Apps bezüglich der Tauglichkeit für alltägliche Gespräche im Gesundheitswesen (Panayiotou et al. 2019)	57
Abbildung 20: Individualisierbarkeit der Spracheingabe bzw. -ausgabe; Vielzahl an Dialekten (Screenshot)	68
Abbildung 21: Nonverbale Kommunikation des Arztes während der Untersuchung (Ausschnitt aus dem Video)	105
Abbildung 22: Körpersprache während des Experimentes; Händeschütteln bei der Begrüßung; Ende der Untersuchung; Patient nickt, Arzt hält Daumen nach oben (Ausschnitt aus dem Video).....	107
Abbildung 23: Der Arzt verwendet eine Grafik, um die Krankheit an den Beinvenen zu erklären (Ausschnitt aus dem Video)	107
Abbildung 24: Manuelle Korrekturfunktion; der Patient zeigt dem Arzt, wie man die Transkription bearbeiten kann; der Arzt korrigiert den Text mit Hilfe einer Tastatur (Ausschnitt aus dem Video).....	109

Einleitung

In unserer hochtechnisierten Welt ist das Leben der Menschen stark von digitalen Technologien beeinflusst. Die neuen Medien stellen in vielen Bereichen eine große Erleichterung dar und können sinnvoll eingesetzt werden, um viele alltägliche Herausforderungen meistern zu können. Stellen wir uns vor, wir reisen ins Ausland und befinden uns in einer Stadt, in der wir uns nicht auskennen. Zudem sprechen wir die Landessprache nicht. Möchten wir den schnellstmöglichen Weg zum Hotel finden, öffnen wir eine Navigations-Applikation auf dem Smartphone, die uns problemlos zum Ziel führt. Sie informiert uns möglicherweise sogar darüber, wo es Radar- oder Verkehrskontrollen gibt, welche Geschwindigkeitsbeschränkungen es gibt und wo wir tanken und parken können. Auf der Suche nach einem guten Restaurant können wir entweder in Google „das beste Restaurant in Rom“ eingeben, oder, wenn wir kommunikationsfreudiger sind und auf die Meinung der EinwohnerInnen Wert legen, auch jemanden auf der Straße fragen. Natürlich sind Sprachbarrieren auch kein Problem mehr: Mit Hilfe von online Dolmetsch-Applikationen (engl. speech-to-speech translation applications, SST) können wir uns leicht verständigen und in ein paar Minuten den gewünschten Ratschlag bekommen. Das Navigationssystem führt uns wieder zum Ziel und nach kurzer Zeit genießen wir die köstlichsten Speisen der Stadt. Diese Geschichte scheint zwar etwas idealisiert zu sein, ist aber durchaus nicht fern der Realität.

Doch auch wenn derartige Applikationen und Programme rasant weiterentwickelt werden, ist die Frage berechtigt, ob SST-Applikationen in unterschiedlichen Einsatzgebieten des täglichen Lebens wirklich so gut und sicher verwendbar sind, wie die Werbung verspricht. Wenn ja, warum gibt es dann weiterhin eine Dolmetschausbildung und diplomierte DolmetscherInnen? Um diese Frage zu beantworten, muss die Komplexität des Dolmetscheinsatzes näher beobachtet werden. Im erwähnten Fall gibt es eine einfache Frage, auf die wahrscheinlich eine relativ schnelle und eindeutige Antwort gegeben wird. Das Gespräch hat auch keine schwerwiegenden Folgen. Würde man sich nicht klar verstehen, wäre die Google-Suche weiterhin eine gute Option. Anders verhält es sich aber, wenn SST-Applikationen in Situationen verwendet werden, die durch eine größere Komplexität und eine Fachsprache gekennzeichnet sind, wie beispielsweise im Gesundheitsbereich. Missverständnisse oder Fehlinterpretationen könnten hier nicht nur gravierende gesundheitliche, sondern auch rechtliche Konsequenzen haben. Einerseits sind es Slogans wie „iTranslate – Your Passport to the World“ oder „SayHi – The Interpreter In Your Pocket“, die uns ein falsches und idealisiertes Bild über diese modernen

Systeme liefern, andererseits ist es der Mangel an wissenschaftlichen Bewertungen und Evaluierungen solcher Applikationen in realitätsnahen Situationen.

Ziel meiner Masterarbeit ist daher eine aufgabenbezogene Evaluierung einer der bekanntesten Dolmetsch-Applikationen: *SayHi Translate*. Die App wird im Rahmen eines Arzt-Patienten-Gespräches im Sprachenpaar Deutsch-Ungarisch getestet und nach verschiedenen Aspekten evaluiert. Es wird untersucht, inwiefern *SayHi Translate* dafür geeignet ist, ein Arzt-Patient-Gespräch zwischen Deutsch und Ungarisch zu dolmetschen.

Um einen Überblick über die Entwicklung des maschinellen Dolmetschens im Laufe der Zeit zu geben, werden nach einer kurzen terminologischen Klärung die wichtigsten Meilensteine der Geschichte des maschinellen Übersetzens und die Vorläufer der maschinellen Dolmetsch-Systeme dargestellt. Danach wird die Funktionsweise der Dolmetsch-Systeme erläutert, indem die verschiedenen Bestandteile (automatische Spracherkennung, maschinelle Übersetzung und Sprachsynthese) beschrieben werden. Danach wird die Applikation *SayHi Translate* vorgestellt. Um die Problematik des DolmetscherInnenmangels im Gesundheitswesen zu verstehen, wird zunächst die Situation in Österreich beleuchtet und die komplexe Rolle der DolmetscherInnen im Gesundheitswesen erklärt. Die Untersuchungen und Forschungsarbeiten, die in diesem Bereich bereits durchgeführt wurden, sowie Ergebnisse und Schlussfolgerungen aus vorhandenen Masterarbeiten werden in Kapitel 6 ausgeführt. Nach der Experimentbeschreibung wird die Evaluierungsmethode geschildert und die Analyse anhand des Experimentes durchgeführt. Schließlich werden die Ergebnisse der Evaluierung zusammengefasst und die Schlussfolgerungen formuliert.

1. Begriffserklärung

Die vorliegende Arbeit setzt sich mit dem Thema „maschinelles Dolmetschen“ auseinander. Um diesen Begriff erläutern zu können, muss zuerst der Unterschied zwischen „Übersetzen“ und „Dolmetschen“ dargestellt werden, die in der Umgangssprache oft verwechselt und als Synonyme verwendet werden.

Unter Übersetzen wird die schriftliche Übertragung eines Ausgangstextes (AT) in einen anderssprachigen Zieltext (ZT) verstanden, während sich Dolmetschen auf die mündliche Übertragung von Informationen in eine andere Sprache bezieht (vgl. Schmitz 1998: 6). Beim Übersetzen besteht jederzeit die Möglichkeit, den Zieltext zu ändern oder zu korrigieren, wohingegen dies beim Dolmetschen kaum möglich ist, wie es auch in den Definitionen von Otto Kade (1968: 35) beschrieben wird:

Wir verstehen [...] unter Übersetzen die Translation eines fixierten und demzufolge permanent dargebotenen bzw. beliebig oft wiederholbaren Textes der Ausgangssprache in einen jederzeit kontrollierbaren und wiederholt korrigierbaren Text der Zielsprache.

Unter Dolmetschen verstehen wir die Translation eines einmalig (in der Regel mündlich) dargebotenen Textes der Ausgangssprache in einen nur bedingt kontrollierbaren Text der Zielsprache.

Allerdings darf auch die Tätigkeit des Gebärdensprachdolmetschens nicht vergessen werden, wobei das Dolmetschen hier nicht nur auf Basis eines gesprochenen Ausgangstexts erfolgt (vgl. Pöchhacker 2016: 10). Es kann auch vorkommen, dass sich Übersetzen und Dolmetschen miteinander vermischen, wie es auch bei Live-Untertitelungen von Fernsehsendungen der Fall ist, wo der ursprünglich mündlich produzierte Input in kürzester Zeit in einen schriftlichen Zieltext übertragen werden muss (vgl. Jüngst 2010: 138).

Wie die Tätigkeit des Übersetzens je nach Fachrichtung weiter unterteilt werden kann (z.B. Literaturübersetzen, Fachübersetzen usw.), so kann auch zwischen verschiedenen „Formen“ des Dolmetschens unterschieden werden. Grundsätzlich kann das Dolmetschen – je nachdem, ob die Verdolmetschung während oder erst nach der ausgangssprachlichen Rede erfolgt – in zwei Dolmetschmodi gegliedert werden, nämlich Simultan- und Konsekutivdolmetschen. Im Simultanmodus wird gleichzeitig, während der AT vorgetragen wird, gedolmetscht; beim Konsekutivmodus erst, nachdem die Ausgangsrede beendet wurde (vgl. Pöchhacker 2016: 18). Zu weiteren Modi des Dolmetschens zählen „Vom-Blatt-Dolmetschen“, bei dem DolmetscherInnen einen schriftlich vorliegenden AT mündlich in die Zielsprache übertragen;

Flüsterdolmetschen (Chuchotage) und Dialogdolmetschen, wobei Letzteres sowohl simultan als auch konsekutiv erfolgen kann (vgl. Prunč 2011: 23f).

Um Mitte des 20. Jahrhunderts, nach der Entwicklung der ersten Computer, etablierte sich der Begriff der maschinellen Übersetzung (MÜ) (vgl. Ramlow 2009: 17). Durch die technologischen Fortschritte im Bereich der MÜ entwickelte sich ab den 1990er Jahren der Bereich des maschinellen Dolmetschens (MD), indem die MÜ-Komponente durch zwei weitere Bestandteile, die automatische Spracherkennung und die Sprachsynthese, ergänzt wurde (vgl. Waibel 2015: 104, 108). Wie diese Entwicklung im Laufe der Jahrzehnte zustande kam, wird im folgenden Kapitel dargestellt.

2. Geschichte des maschinellen Dolmetschens

Die drei Komponenten der maschinellen Dolmetsch-Systeme wurden über die Jahrzehnte unabhängig voneinander entwickelt. Im Bereich der maschinellen Übersetzung wurden bereits ab den 1950er Jahren Forschungsarbeiten erstellt, wohingegen die Systeme für automatische Spracherkennung (engl. automatic speech recognition, ASR) und für Sprachsynthese (engl. text-to-speech synthesis, TSS) erst später entwickelt wurden (vgl. Waibel & Fügen 2008: 73). In den folgenden Kapiteln werden die Meilensteine in der Entwicklung der maschinellen Übersetzungssysteme im Laufe der Geschichte dargelegt bzw. erläutert, wie die ersten MD-Systeme – dank der technologischen Fortschritte im Bereich der MÜ – entstanden sind.

2.1 Anfänge der maschinellen Übersetzung

Der Traum, den Übersetzungsprozess zu automatisieren, also eine Maschine zur Verfügung zu haben, die es ermöglicht, problemlos Sprachbarrieren zu überwinden, existiert schon seit dem 17. Jahrhundert (vgl. Ramlow 2009: 11). Die ersten Versuche, diesen Traum zu verwirklichen, wurden allerdings erst im Jahre 1933 unternommen, als der französische Ingenieur Georges Artsrouni und der Russe Petr Petrovich Smirnov-Trojanskij voneinander unabhängig maschinengestützte Übersetzungssysteme entwickelten. Diese Systeme konnten nur die Grundformen der Wörter berücksichtigen und beruhten auf einem maschinellen Vergleich verschiedener Lexika. Deshalb war eine Vor- und Nachedition durch einen Menschen beim Übersetzungsprozess stets erforderlich (vgl. Schwanke 1991: 69).

Offiziell hat sich der Begriff der „maschinellen Übersetzung“ erst nach der Erfindung des Computers (1942) im Jahre 1949 etabliert, als das sogenannte „Weaver Memorandum“ von Warren Weaver veröffentlicht wurde. Weaver, der in der Literatur oft als Vater der maschinellen Übersetzung bezeichnet wird, würdigte in seinem Memorandum die bisherigen Forschungsergebnisse im Bereich der MÜ und vertrat die optimistische Auffassung, dass die offensichtlichen Herausforderungen in der automatischen Übersetzung – vor allem das Problem der Mehrdeutigkeit der Zeichen – prinzipiell lösbar seien, indem die Sprache als eine Art Code behandelt werde. Er stützte seine optimistischen Gedanken auf die Vorstellung, dass alle Sprachen über eine logische Struktur verfügen und somit auch die sprachlichen Universalien immer gegeben seien. Er vertrat die Ansicht, dass es grundsätzlich möglich sei, sogar unbekannte Zeichen – wie die feindlichen Korrespondenzen im Zweiten Weltkrieg – mit Hilfe von Dekodierung zu identifizieren und zu übersetzen (vgl. Schwanke 1991: 70; Hutchins 2015:

120). Als Folge des Weaver Memorandums wurde die Forschung im Bereich der MÜ gefördert. Es wurden zahlreiche Forschungsgruppen gebildet, unter anderem die Arbeitsgruppen am *Massachusetts Institute of Technology (MIT)* und an der *University of California at Los Angeles (UCLA)* (vgl. Schwanke 1991: 71). Die Welle an Forschungsarbeiten führte dazu, dass bereits im Jahr 1954 das erste MÜ-System präsentiert wurde, das – offensichtlich für politische Zwecke – eine vollautomatische Übersetzung vom Russischen ins Englische lieferte. Das System verfügte mit 250 Lexikon-Einträgen über einen beschränkten Wortschatz und berücksichtigte bei der Übersetzung lediglich sechs grammatische Regeln, es wurde in der damaligen Zeit jedoch als bahnbrechender Erfolg angesehen (vgl. Schwanke 1991: 72).

2.2 Der ALPAC-Bericht

In den darauffolgenden Jahren wurden zahlreiche Forschungsarbeiten im Bereich der MÜ erstellt; in Fachkreisen herrschte großer Optimismus bezüglich der Möglichkeiten der Qualitätsverbesserung von maschinellen Übersetzungen. Dieser Aufschwung ging jedoch im Jahr 1966 plötzlich zu Ende, als der sogenannte ALPAC-Bericht (*Automatic Language Processing Advisory Committee*) veröffentlicht wurde. Im ALPAC-Bericht wurde der aktuelle Forschungsstand im Bereich der MÜ beschrieben und darauf hingewiesen, dass zu jener Zeit kein Übersetzungssystem in der Lage sei, eine Übersetzung zu leisten, die auch ohne Verbesserung zufriedenstellend wäre. Es wurde außerdem festgestellt, dass keine Maschine Texte generieren kann, die im Hinblick auf Informationsgehalt und Verständlichkeit die Qualität von HumanübersetzerInnen erreichen könne. Die Notwendigkeit einer Postedition durch menschliche ÜbersetzerInnen wurde betont und im Hinblick auf die Bemühungen, eine vollautomatische maschinelle Übersetzung zu generieren, als Misserfolg wahrgenommen (vgl. Hutchins 1986: 165f). Die Studien und Evaluierungen – auf die sich der Bericht stützte – zeigten, dass maschinelle Übersetzungen zwar teilweise verständlich waren, aber oft falsch übersetzte und irreführende Informationen aufwiesen (vgl. ALPAC 1966: 19). Der ALPAC-Bericht führte zu einem erheblichen Imageverlust der MÜ, auch die finanziellen Mittel für die weitere Forschungsarbeiten wurden drastisch gekürzt (vgl. Hutchins 1986: 167).

2.3 Das Wiedererwachen der Forschung

Die 1970er und 80er Jahre waren durch Ernüchterung gekennzeichnet. Es wurde versucht, realistischere Ziele bezüglich der Leistung der MÜ-Systeme zu setzen. Die Zielsetzung der Forschung war – anstatt der Generierung hochqualitativer Übersetzungen – nun die Entwicklung

von Systemen, welche die Inhalte des Ausgangstextes in einer verständlichen Form wiedergeben können (vgl. Hutchins 1986: 175).

Bereits im Jahr 1968 konnten Fortschritte im Bereich der MÜ verzeichnet werden, als das maschinelle Übersetzungssystem *Systran* entwickelt wurde, das vor allem in den Sprachenpaaren *Englisch – Russisch* bzw. *Englisch – Französisch* Übersetzungen lieferte. Es wurde ab 1970 von der US Air Force eingesetzt und seit 1974 auch von der NASA verwendet. *Systran* basierte auf einem direkten Übersetzungsansatz (Kapitel 3.2.1.1), genauso wie das Übersetzungssystem *Logos* (1969), das in erster Linie zur Übersetzung von Flugzeughandbüchern vom Englischen ins Vietnamesische eingesetzt (1970-1972), später aber durch weitere Sprachenpaare ergänzt wurde (vgl. Koehn 2010: 16; Ramlow 2009: 73).

Ab Mitte der 70er Jahre wurde die Künstliche Intelligenz ein wesentlicher Faktor für die Weiterentwicklung maschineller Übersetzungssysteme. Von der Europäischen Gemeinschaft wurde das Forschungsprojekt *EUROTRA* ins Leben gerufen, wobei lexikalische, syntaktische und semantische Informationen in den Übersetzungsprozess einbezogen wurden, nicht aber das außersprachliche Wissen. Obwohl das Projekt in der Industrie kaum auf Interesse stieß und kein einsatzfähiger Prototyp entwickelt wurde, hat *EUROTRA* europaweit neue Impulse für künftige Forschungsprojekte gegeben (vgl. Hutchins 1995: 438f).

In den 1980er Jahren wurde auch der Versuch unternommen, indirekte Übersetzungssysteme zu entwickeln, die entweder interlingua- oder transferbasiert arbeiteten (Kapitel 3.2.1). Zu solchen Systemen zählten z.B. das experimentelle Projekt *EURATOM*, das japanische Projekt *CICC*, das System *Mu* in Kyoto, das *GETA-Ariane System* in Grenoble sowie das wissensbasierte Übersetzungsprojekt der Carnegie-Mellon University in Pittsburgh (vgl. Hutchins 2015: 123f).

2.4 Die Entwicklung maschineller Dolmetsch-Systeme

Anfang der 1990er Jahre – als die ersten ARS-, MT- und TSS-Systeme einen erwünschten Reifegrad erreichten – wurde versucht, diese drei zu integrieren. So entstanden die ersten maschinellen Dolmetsch-Systeme, die über die folgenden drei Jahrzehnte ständig weiterentwickelt und verbessert wurden. Dank der enormen technologischen Fortschritte sind die heutigen MD-Systeme in Form von Dolmetsch-Applikationen auf Smartphones und Tablets anwendbar und verfügen über eine große Anzahl verschiedener Sprachenpaare. Die wichtigsten Projekte und Systeme im Bereich des maschinellen Dolmetschens werden in Tabelle 1 zusammenfassend dargestellt:

Tabelle 1: Entwicklungsphasen maschineller Dolmetsch-Systeme (Waibel 2015: 107)

	Jahre	Vokabular	Sprechstil	Domäne	Geschwindigkeit	Plattform	Beispielsysteme
Erste Dialog-Demonstrationssysteme	1989–1993	begrenzt	beschränkt	limitiert	$2-10 \times \text{RT}$	Workstation	JANUS-1, C-STAR-I
Einfache Phrasenbücher	seit 1997	begrenzt, modifizierbar	beschränkt	limitiert	$1-3 \times \text{RT}$	Handheld-Gerät	Phraselator, Ectaco
Spontane Zwei-Weg-Systeme	seit 1993	unbegrenzt	spontan	limitiert	$1-5 \times \text{RT}$	PC/Handheld-Gerät	JANUS-III, C-STAR, Verbmobil, Nespole, Babylon, Transtac
Übersetzung von Nachrichtensendungen, politischen Reden	seit 2003	unbegrenzt	Lesen/Vorbereitete Rede	offen	offline	PCs, PC-Clusters	NSF-STRDUST, EC TC-STAR, DARPA GALE
Simultanübersetzen von Vorträgen	seit 2005	unbegrenzt	spontan	offen	Echtzeit	PC, Laptop	Lecture Translator
Kommerzielle Konsekutivübersetzer per Telefon	seit 2009	unbegrenzt	spontan	offen	online und offline	Smartphone	Jibbiggo, Google, Microsoft
Simultan-dolmetscher-Leistungen	seit 2012	unbegrenzt	spontan	offen	online	Server, Cloud-basiert	KIT, EU-Bridge, Microsoft

Wie Tabelle 1 zeigt, waren die ersten Modelle der MD-Systeme, die zwischen 1989 und 1993 entwickelt wurden, noch sehr limitiert einsatzfähig. Sie verfügten nur über ein kleines Vokabular (weniger als 1000 Wörter), stellten hohe Anforderungen an Genauigkeit und Geschwindigkeit des eingesprochenen Textes und beschränkten sich auf limitierte Gesprächsdomänen und Anwendungsszenarien (vgl. Waibel 2015: 108). Zu den ersten Demonstrationssystemen zählte unter anderem das *JANUS*-System, das 1991 der Öffentlichkeit vorgestellt wurde. Es entstand im Rahmen einer Zusammenarbeit der Universität Karlsruhe und der Carnegie-Mellon-Universität in Pittsburgh mit den Forschungsgruppen der *ATR Interpreting Telephony Laboratories* in Japan, die parallel an der Entwicklung ähnlicher Systeme für die japanische Sprache arbeiteten. Das Dolmetsch-System *JANUS* wurde ab 1993 für Telefongespräche über die Registrierung für Konferenzen mit den Sprachen Deutsch, Japanisch und Englisch eingesetzt (vgl. Hutchins 1995: 442f; Waibel 2015: 108; Waibel et al. 1991).

Da diese eingeschränkten Modelle nur für die Verarbeitung vorgeschriebener Dialoge geeignet und deshalb für spontansprachliche Einsätze unbrauchbar waren, war das nächste Ziel die Entwicklung eines MD-Systems, das auch unter der Verwendung von Spontansprache benutzt werden kann. So entstanden ab 1993 Systeme, die zwar in der Lage waren, einige spontansprachliche Merkmale (wie Füllwörter, Wiederholungen, Pausen usw.) zu erkennen und herauszufiltern, aber weiterhin nur für einen vorab definierten, beschränkten Themenbereich anwendbar waren. In diesem Zeitraum (1993-2000) wurde ein enormer Forschungsaufwand betrieben und zahlreiche Systeme wurden vorgestellt, die auch mit spontan gesprochenen Texteingaben arbeiten konnten, unter anderem die *C-STAR-Systeme* (*Consortium for Speech Translation Advanced Research*), *NESPOLE*, *JANUS III*, *BABYLON*, *TRANSTAC* und *VERBMOBIL* (vgl. Waibel 2015: 109; Waibel & Fügen 2008: 74).

VERBMOBIL¹ war ein langfristig angelegtes Projekt, das 1993 in Deutschland ins Leben gerufen wurde und als Meilenstein in der Geschichte des maschinellen Dolmetschens betrachtet werden kann. Das Projekt wurde vom deutschen Bundesministerium für Bildung, Wissenschaften und Forschung (BMBF) gefördert und hatte das Ziel, ein tragbares Dolmetschsystem zu entwickeln, das spontansprachliche Dialoge vor allem aus den Domänen Reiseplanung, Hotelreservierung und Fernwartung in der Sprachkombination Deutsch-Englisch-Japanisch „dolmetschen“ kann (vgl. Wahlster 2000: 20ff). Das VERBMOBIL-System wurde zwischen 1993 und 2000 in zwei Phasen entwickelt (1993-1996, 1997-2000) und kann als äußerst innovativ angesehen werden, da es die verschiedenen Bedeutungsvarianten aus den unterschiedlichen Betonungen der Wörter erkannte bzw. in der Lage war, eine Gesprächsgeschichte aus den Dialogen zu erstellen (vgl. 2000: 22, 119, 338). Das System war das erste „speech-only translation system“. Das bedeutet, dass die Nutzung des Systems auch ohne Tastenwahl, rein durch Sprachbefehle möglich war (vgl. 2000: 4).

Da diese Systeme weiterhin schwere und große Hardware benötigten, war der nächste Schritt die Entwicklung tragbarer Dolmetsch-Systeme für PDAs, Smartphones und Laptops. Wegen der schlechten Hardware-Voraussetzungen und geringen Speicherkapazitäten der damaligen Geräte war dies jedoch eine sehr große technische Herausforderung (vgl. Waibel & Fügen 2008: 74).

Zu den ersten tragbaren Systemen zählten unter anderem *Phraselator* und *pocket translator*. Beide Systeme konnten nur für beschränkte Themenbereiche (wie z.B. Reisen, Medizin

¹ „VERB“ – **verb**ale Kommunikation, „MOBIL“ – in **mob**ilen Situationen (vgl. Wahlster 2000: 3)

oder Militär) und unter Anwendung eines korrekten Sprachstils verwendet werden (vgl. Waibel & Fügen 2008: 74).

Ab 2003 wurde versucht, das Problem der Domänenabhängigkeit zu lösen, indem Dolmetsch-Systeme entwickelt wurden, die sowohl mit Spontansprache arbeiten als auch in jeglichen Kontexten Anwendung finden konnten. Dies erwies sich jedoch als eine sehr komplexe und herausfordernde Aufgabe, für die mehrere Forschungsprojekte ins Leben gerufen wurden. Die wichtigsten Projekte waren *STR-DUST* (*Speech Translation for Domain-Unlimited Spontaneous Communication Tasks*), *TC-STAR* (*Technology and Corpora for Speech to Speech Translation*) sowie *DARPA GALE* (*Global Autonomous Language Exploitation*) aus den Jahren 2003, 2004 und 2006 (vgl. Waibel & Fügen 2008: 75). Eines der ersten Systeme, das in domänenunabhängigen Settings im Simultanmodus Inhalte übertragen konnte, war der *Lecture Translator*, dessen Entwicklung im Jahr 2005 begann. Ziel dieses MD-Systems war es, eine zeitgleiche maschinelle Verdolmetschung von Vorlesungen für Studierende aus dem Deutschen ins Englische zu ermöglichen. Die Leistung von *Lecture Translator* war jedoch noch nicht optimal, da das System während der Vorlesungen mit verschiedenen Herausforderungen konfrontiert wurde: Einerseits musste die ASR-Komponente die Disfluenzen (wie z.B. Selbstkorrekturen, Häitationen, Wortwiederholungen) der spontan vorgetragenen Rede erkennen und korrigieren, andererseits kamen während der Präsentationen sehr oft Fachbegriffe vor, welche die Arbeit der MÜ-Komponente erschwerten (vgl. 2008: 76).

Dank der zunehmenden Verbreitung von verschiedenen mobilen Endgeräten sowie der enormen technologischen Fortschritte (Breitband-Internetzugang, Künstliche Intelligenz, Machine Learning usw.) sind heutzutage unzählige Dolmetsch-Applikationen auf dem Markt zu finden, die nicht nur eine Leistungsverbesserung aufweisen, sondern auch einfach auf Smartphones oder Tablets heruntergeladen oder im Browser aufgerufen werden können. Zu den meist verwendeten Dolmetsch-Apps gehören unter anderem *Google Translate*, *Skype Translator*, *iTranslate*, *Microsoft Translator*, *MyLingo*, *TripLingo* und *SayHi Translate*. In der vorliegenden Masterarbeit wird *SayHi Translate* im Rahmen einer ärztlichen Untersuchung getestet und anschließend anhand verschiedener Aspekte evaluiert.

Da nun die geschichtliche Entwicklung des maschinellen Dolmetschens dargestellt wurde, wird in den nächsten Kapiteln die Funktionsweise der MD-Systeme beschrieben.

3. Funktionsweise maschineller Dolmetsch-Systeme

Wie bereits erwähnt, setzen sich MD-Systeme grundsätzlich aus drei Hauptkomponenten zusammen: aus der automatischen Spracherkennung, der maschinellen Übersetzung und der Sprachsynthese. In diesem Kapitel werden diese Bestandteile genauer vorgestellt.

3.1 Automatische Spracherkennung

Die erste Komponente maschineller Dolmetschsysteme ist die automatische Spracherkennung. Obwohl BenutzerInnen den Eindruck bekommen können, dass die Maschine intelligent genug wäre, um sie zu verstehen, handelt es sich in Wirklichkeit nur um Algorithmen, die die Spracherkennung ermöglichen. Sogar die besten heute existierenden Systeme sind noch weit von einem echten Verständnis entfernt. Sie sind aber in der Lage, menschliche Sprache aus einem Audiosignal in Text zu konvertieren – seit kurzem auch die natürliche, fließend gesprochene Sprache (vgl. Kolossa 2017: 2).

Spracherkennungssysteme können je nach Einsatzgebiet in verschiedene Kategorien eingeteilt werden. Die grundlegende und einfachste Form der Spracherkennung ist die Einzelworterkennung, wobei das Programm bei jedem Schritt ein isoliert gesprochenes Wort erkennen und digitalisieren muss. Dieses Konzept eignet sich am besten für Anwendungen, die durch Kommandowörter (z.B. „Pause“, „Lauter“, „Antworten“) steuerbar sind (vgl. Bub 1999: 4).

In einer weiteren Kategorie kann man Systeme mit Wortkettenerkennungsfunktion zusammenfassen, also Programme, die in der Lage sind, Wörter, die ohne eine besondere Grammatik aneinandergereiht sind, zu erkennen. Diese Systeme können bei der Erkennung von Ziffernketten, wie zum Beispiel Kreditkartennummern oder Telefonnummern, Anwendung finden (vgl. Bub 1999: 4).

Eine weitere Kategorie beinhaltet Programme mit kontinuierlicher Spracherkennung. Diese Systeme können ganze Sätze verarbeiten, da sie in der Lage sind, fließend aneinander gekettete Wörter zu erkennen, deren Abfolge durch Grammatikregeln bestimmt wird. In diese heutzutage bedeutendste Kategorie fallen jene Spracherkennungssysteme, die für die Erkennung spontan geäußerter Sprache konzipiert sind. Hierbei erkennt das Programm die Sprache, die während der natürlichen Kommunikation verwendet wird, mit all ihren zahlreichen Füllwörtern und Wortwiederholungen (vgl. Bub 1999: 4f). Diese letzte Kategorie ist aus Sicht dieser Arbeit am wichtigsten, da Dolmetsch-Applikationen meist in spontanen Gesprächen

verwendet werden und für die Erkennung längerer Sinneinheiten geeignet sind. Wie genau die automatische Spracherkennung in solchen Systemen funktioniert, wird im Folgenden beleuchtet.

Der wichtigste Gegenstand, um die lautsprachlichen Informationen einer gesprochenen Äußerung auswerten und transkribieren zu können, ist das Sprachsignal, das entsteht, indem die von einer Person produzierten Schallwellen über einen elektroakustischen Wandler (Mikrofon) in ein elektrisches Signal umgewandelt werden (vgl. Pfister & Kaufmann 2017: 25).

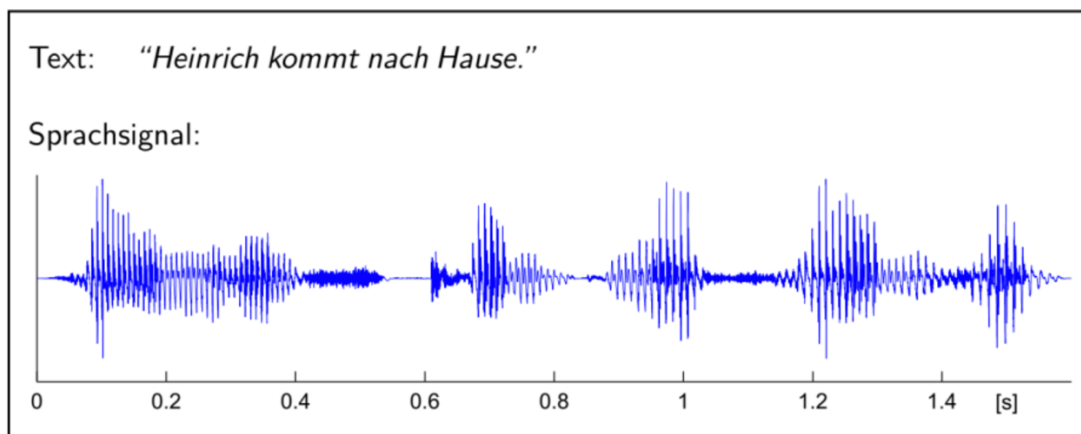


Abbildung 1: Veranschaulichung der Verschiedenheit von Text und Sprachsignal (Pfister & Kaufmann 2017: 7)

Nach der Auswertung durch den Algorithmus wird das Ergebnis als Folge von Wörtern dargestellt. Da sich aber in der gesprochenen Sprache viele Mehrdeutigkeiten verstecken (der Satz „*This machine can recognize speech*“ wird fast genauso ausgesprochen wie „*This machine can wreck a nice beach*“), benötigt die Spracherkennung die Fähigkeit zu erkennen, welche von mehreren ähnlichen Alternativen in einem gegebenen Kontext die sinnvollere und logischere Interpretation des Gesagten ist. Um dies zu erreichen, werden in modernen Spracherkennungssystemen akustische Modelle verwendet, die jedem Laut eine Wahrscheinlichkeit zuordnen. Zusätzlich sind diese Systeme auch mit Aussprachewörterbüchern bzw. mit Sprachmodellen, die die Wahrscheinlichkeit jeder möglichen Wortsequenz des Satzes auswerten, ausgestattet (vgl. Waibel 2015: 105).

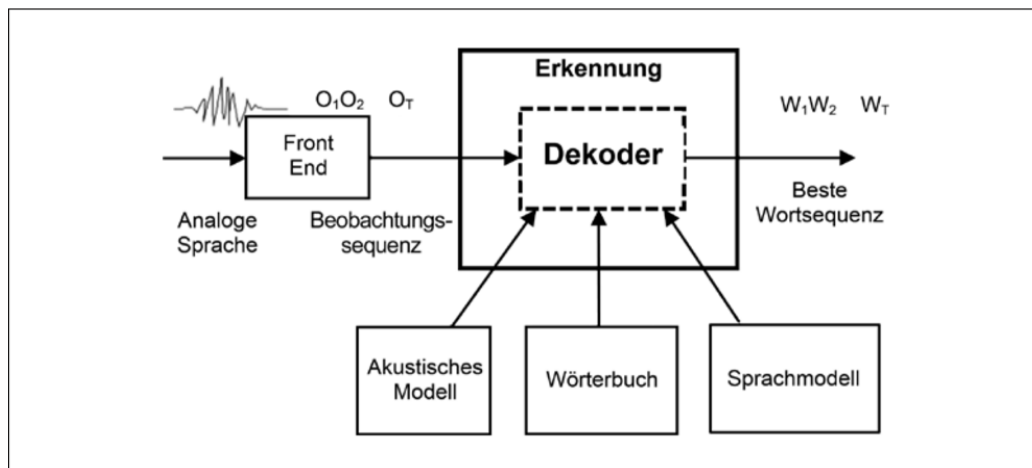


Abbildung 2: Ein typischer Spracherkenner (Speech-to-Text) (Waibel 2015: 105)

Die Auswertung dieser Modelle während der Spracherkennung erfolgt durch automatische Such- und Optimierungsalgorithmen. Um die bestmögliche Transkription zu erzeugen, benutzen moderne Systeme verschiedene Lernalgorithmen, welche ihre Leistung fortdauernd verbessern. Dazu werden neben Optimierungsmethoden auch Neuronale Netze verwendet, die in der Lage sind, anhand von bekannten Beispieldaten die bestmögliche Lösung herauszufiltern (vgl. Waibel 2015: 105f). Dank der konsequenten Weiterentwicklung von komplexen Algorithmen und Neuronalen Netzen konnte man im Forschungsgebiet der automatischen Spracherkennung in den vergangenen Jahren bedeutende Fortschritte machen und die Leistung der Systeme deutlich verbessern.

3.1.1 Herausforderungen der automatischen Spracherkennung bei spontan gesprochener Sprache

Die meisten MD-Systeme arbeiten mit Spracherkennungssystemen, die die spontane, „natürliche“ Sprache der RednerInnen in ein digitales Sprachsignal umwandeln. Man kann das Sprechen im Sinne eines Prozesses auffassen, wobei die sprachliche Aussage die Eingabe des Sprechprozesses darstellt und das Sprachsignal die Ausgabe. Man darf aber nicht außer Acht lassen, dass nicht nur die Aussage den Sprachprozess beeinflusst, sondern auch die Stimme der Sprechenden Person. Sprechgewohnheiten, Dialekte, Akzente, aber auch die Physiologie des Vokaltraktes können einen starken Einfluss auf die Spracherkennung haben und eine große Herausforderung für das System darstellen. Ein weiterer Störfaktor ist die starke Situationsabhängigkeit der Aussprache; der jeweilige emotionale Zustand und eventuell sogar der Gesundheitszustand können sich im Sprachsignal niederschlagen und die Spracherkennung

beeinflussen, wie auch in Abbildung 3 veranschaulicht wird (vgl. Pfister & Kaufmann 2017: 25). Schönert (2012) beschreibt diese Problematik folgenderweise:

Das Problem bei der elektronischen Spracherkennung ist, dass niemand einen Begriff in jeder Situation immer ganz genau gleich ausspricht. Mal ist er müde, mal hektisch, mal laut, mal leise, mal konzentriert, mal betrunken, mal sauer, mal verliebt, mal erkältet. Deshalb ist es für eine Software sehr schwierig, Wörter durch Suchen deckungsgleicher Tonfolgen zu erkennen. (Schönert 2012: 2)

Auch Umgebungsgeräusche üben einen Einfluss auf den Sprechprozess aus, z.B. weil eine Person in einem lauten Umfeld unweigerlich lauter spricht. Zudem können die Umgebungsgeräusche an sich die Signalübertragung beeinflussen, da sie ja selbst Schallwellen produzieren. Die Raumakustik (Übertragung der Schallwellen), die Mikrophoncharakteristik und die Signalkodierung und -kompression (Übertragung des elektrischen Signals) spielen ebenfalls eine wesentliche Rolle in diesem Prozess (vgl. Pfister & Kaufmann 2017: 25).

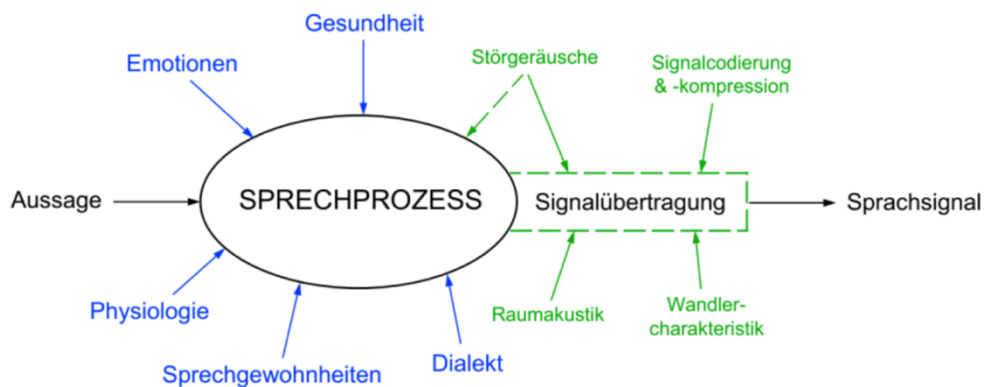


Abbildung 3: Einflussfaktoren auf automatische Spracherkennungssysteme nach Pfister & Kaufmann (2017: 25)

Nicht zu vergessen sind auch die Disfluenzen von spontan gesprochenen Äußerungen, die neben den bereits erwähnten Störfaktoren die Bearbeitung einer Aussage deutlich erschweren können. Unter Disfluenzen versteht man die Unterbrechungen des normalen Sprechflusses, also z.B. Selbstkorrekturen („In äh auf welchem Sender ZDF oder ARD“), Wortwiederholungen („und zwischendurch könnten wir mal von Fernsehen auf auf Radio umschalten“), Abbrüche („wie häufig hat jetzt eigentlich...wie häufig ist der Fernseher eigentlich heute schon eingeschaltet worden“), Anakoluthe, d.h. Sätze, deren Ende dem Anfang grammatikalisch nicht entspricht („wenn wir noch keine Nachrichten sehen können bitte ich den Fernseher solange auszuschalten und um zwanzig Uhr wieder einschalten“), Häsitationen („Ich möchte ... äh ...“), Herausstellungen („und den Ventilator können Sie den bitte anschalten“) oder Nachschübe

(„*ich hätte jetzt doch gern die Lampe angeschaltet und zwar den Deckenfluter*“) (vgl. Fünfer 2013: 15; Hacker et al. 2011: 8).

Der Umfang des Vokabulars des Sprechers bzw. der Sprecherin spielt hier ebenso eine wichtige Rolle. Es kann außerdem vorkommen, dass die sprechende Person Ausdrücke aus anderen Sprachen verwendet, die vom System nicht als solche erkannt werden. Auch beim gleichzeitigen Sprechen mehrerer Personen wird das Programm leicht überfordert (vgl. Fünfer 2013: 15).

Zunächst müssen von einem Spracherkennungssystem, das mit spontan gesprochener Sprache arbeitet, auch Satzgrenzen bzw. für die MÜ geeignete Segmente erkannt werden. Hierbei sind die ASR-Systeme auf Indikatoren wie Prosodie, natürliche Sprachpausen oder auf Sprachmodellen basierende statistische Werte angewiesen. Man kann die Arbeit solcher Systeme wesentlich erleichtern, indem man die Satzzeichen (Punkt, Fragezeichen, Ausrufzeichen, Komma) einfach selbst ausspricht. Diese werden dann vom System erkannt und in Zeichen („“, „?“, „/“, „“, „“) umgewandelt. Neben der sinnvollen Trennung der Inhalte muss das System auch in der Lage sein, Eigennamen zu erkennen (vgl. Fünfer 2013: 15).

Ziel solcher Systeme ist es, eine möglichst niedrige Fehlerquote zu erreichen, im Idealfall soll die sogenannte Wortfehlerrate (word error rate, WER) bei höchstens 10% liegen. Hochwertige und optimal positionierte Mikrophone können die Ausgangsbedingungen deutlich verbessern (vgl. Fünfer 2013: 15).

Auch die Geschwindigkeit des gesamten Prozesses ist ein entscheidender Faktor. Um menschliche DolmetscherInnen nachzuahmen, muss das System schnell und einwandfrei funktionieren. Die optimale Leistung wird erreicht, wenn die Worterkennung möglichst in Echtzeit erfolgt. Bedeutend ist aber nicht nur die Laufzeit des Systems (die Zeit zur Erkennung), sondern auch die zur Verarbeitung eines Segments benötigte Zeit. Analog dem Humandolmetschen darf der Timelag nicht zu groß sein, vor allem wenn ein Sprecherwechsel erfolgt. Um den Prozess zu beschleunigen, werden vom System sogenannte Worthypothesengraphen aus den Informationen erstellt bzw. ein Wahrscheinlichkeitswert ermittelt, um der Übersetzungskomponente mögliche Alternativen weiterzugeben (vgl. Fünfer 2013: 15f).

3.2 Maschinelle Übersetzung

In der zweiten Phase des maschinellen Dolmetschens wird der erkannte und transkribierte Ausgangstext durch maschinelle Übersetzung in die Zielsprache (ZS) übertragen. Wie bereits in Kapitel 2.1 erwähnt, wurden erste Versuche mit maschineller Übersetzung von Texten bereits im zweiten Weltkrieg unternommen. Im Laufe der Zeit wurde jedoch immer eindeutiger, dass trotz der großen Fortschritte, die in diesem Bereich zu verzeichnen sind, keine Maschine Texte generieren kann, die im Hinblick auf den Informationsgehalt und die Verständlichkeit die Qualität der HumanübersetzerInnen erreichen können (vgl. Hutchins 1986: 165f, Waibel 2015: 106).

Da sich die MÜ an der Schnittschnelle von Übersetzungswissenschaft, Informatik, Linguistik und Computerlinguistik befindet, wurden – dank dieser Interdisziplinarität – viele verschiedene Methoden und Ansätze zur MÜ entwickelt (vgl. Tripathi & Sarkhel 2010; Werthmann & Witt 2014: 86), die grundsätzlich in drei Kategorien unterteilt werden können: regelbasierte, statistische-korpusbasierte und neuronale-algorithmenbasierte Strategien.

Regelbasierte Übersetzungsansätze, die in den 80er Jahren entstanden sind, stützen sich in erster Linie auf grammatische Regeln und mehrsprachige Wörterbücher, die nach umfassender linguistischer Analyse beider Sprachen manuell erstellt wurden (vgl. Werthmann & Witt 2014: 86f).

Korpusbasierte Strategien basieren – wie auch der Name verrät – vor allem auf der Analyse von ein- und mehrsprachigen Korpora, mit deren Hilfe das System die Wahrscheinlichkeit unterschiedlicher Übersetzungsmöglichkeiten aus den Wortmustern und -entsprechungen berechnet und die wahrscheinlichste Übersetzung auswählt. Korpusbasierte Ansätze (auch statistisch basierte/statistische MÜ genannt) wurden in den 1980er und 1990er Jahren entwickelt und wurden bei online-Übersetzungssystemen von großen Firmen wie *Google* oder *Microsoft* angewendet (vgl. Werthmann & Witt 2014: 91, 96).

Dank dem exponentiellen Wachstum der Datenmenge und den Entwicklungen im Bereich Künstliche Intelligenz ist seit 2015 ein Paradigmenwechsel in der MÜ zu beobachten: Es hat sich der Ansatz der neuronalen MÜ etabliert, der auf neuronalen Netzen basiert und sich mit Hilfe von Machine Learning ständig weiterentwickelt (vgl. Lionbridge 2017). Diese Übersetzungsstrategie wird in Kapitel 3.2.3 näher erläutert.

3.2.1 Regelbasierte Ansätze

Die regelbasierten Ansätze (engl. Rule-Based Machine Translation, RBMT) zählen zu den klassischen Ansätzen der MÜ und fanden sich lange in den meisten kommerziellen Systemen wieder. RBMT-Systeme arbeiten in drei Stufen: Zuerst wird die AS auf verschiedenen linguistischen Ebenen analysiert (z.B. morphologische, syntaktische Ebene usw.), danach wird sie je nach MÜ-Strategie in eine mehr oder wenige abstrakte Repräsentation überführt (Transfer), auf die schließlich Übersetzungsregeln angewendet werden, um eine fertige Übersetzung zu generieren. Abgängig davon, wie komplex der regelbasierte Übersetzungsprozess abläuft (also wie viele Sprachenpaare bei der Übersetzung berücksichtigt werden bzw. wie die gewonnenen Informationen aus der AS analysiert und abstrakt repräsentiert werden), werden drei MÜ-Strategien unterschieden: direkte Übersetzung, Transfer und Interlingua. In Vauquois' Dreieck werden diese unterschiedlichen Übersetzungsstrategien in einem Architekturschema wie in Abbildung 4 dargestellt (vgl. Werthmann & Witt 2014: 87f).

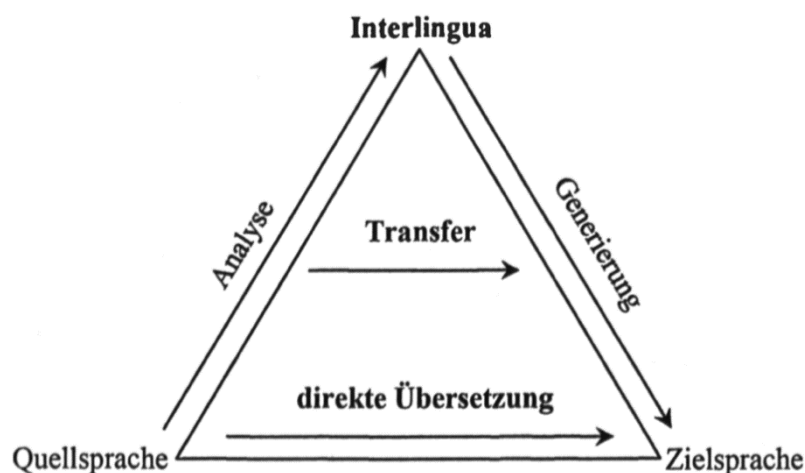


Abbildung 4: MÜ Vauquois' Dreieck der MÜ nach Hutchins & Somers (1992: 107)

Während auf der untersten Ebene des Dreiecks die direkte Übersetzung mit ihrer minimalen Analyse liegt, befindet sich die Interlingua-Übersetzung, die versucht alle Sprachen in einer gemeinsamen sprachunabhängigen Repräsentation abzubilden, an der Spitze des Vauquois-Dreiecks. In der Mitte ist der Transfer-Ansatz zu finden, der eine sprachabhängige Zwischenrepräsentation erzeugt und in diesem Sinn als eine Art „Mittelweg“ betrachtet werden kann. Transfer und Interlingua gehören zu den indirekten Übersetzungsstrategien, die gegenüber der direkten Übersetzung durch eine höhere Komplexität gekennzeichnet sind und auch als

Strategien der 2. Generation bezeichnet werden (vgl. Eberle 2008; Stein 2009; Werthmann & Witt 2014: 87f). Diese drei Ansätze werden im Folgenden genauer vorgestellt.

3.2.1.1 Direkter Ansatz

Die direkte Übersetzung ist der einfachste und älteste Ansatz der MÜ, bei dem das System mit einfachen Wort-zu-Wort-Übersetzungen arbeitet und dabei keine (oder nur eine minimale) strukturelle und semantische Analyse der AS durchführt (vgl. Stein 2009; Werthmann & Witt 2014: 88).

Im Übersetzungsprozess wird zuerst der AT in Wörter bzw. Phrasen segmentiert, die danach mit den Einträgen des bilingualen Wörterbuches bzw. Lexikons verglichen werden, um möglichst genaue Übereinstimmungen zu finden, mit denen die Wörter der AS mit Entsprechungen aus der ZS ersetzt werden können. Da es aber für viele Wörter mehr als eine mögliche Übersetzung gibt, ist die Qualität direkter Übersetzungen nicht besonders hoch und nur in eingeschränkten Anwendungsszenarien zu gebrauchen. Außerdem werden Mehrwortlexeme (z.B. „ins Gras beißen“), die zumeist nicht wörtlich zu übersetzen sind, vom System nicht erkannt. Weiterhin werden die Ergebnisse meist nur oberflächlich an die Satzstellung der ZS angepasst, was sich natürlich auch negativ auf die Übersetzungsqualität auswirkt. Man darf auch nicht unerwähnt lassen, dass Systeme, die mit direkten Übersetzungsstrategien arbeiten, immer nur ein Sprachenpaar berücksichtigen können. Folgende Abbildung zeigt die Funktionsweise dieses einfachen Ansatzes (vgl. Stein 2009; Werthmann & Witt 2014: 88f).

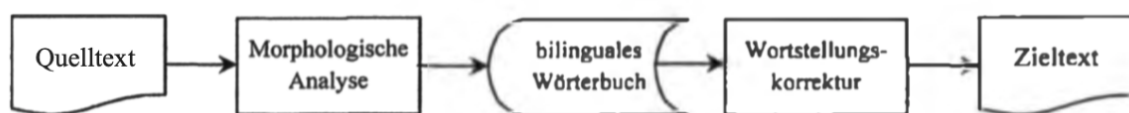


Abbildung 5: Schema für ein direktes MÜ-System (Carstensen et al. 2004: 566)

3.2.1.2 Transfer-Ansatz

Im Gegensatz zur direkten Übersetzung läuft der Übersetzungsprozess beim Transfer-Ansatz nicht auf Wortebene ab, sondern auf Satzebene. Das bedeutet, dass Wörter hier als Teil syntaktischer Strukturen betrachtet werden; nicht einzeln und isoliert, wie bei der direkten Übersetzung (vgl. Werthmann & Witt 2014: 89).

Der AT wird zuerst in Sätze segmentiert und erst danach morphologisch, semantisch und syntaktisch analysiert. Nach der Analyse wird der AT in eine abstrakte Struktur übertragen,

die im nächsten Schritt in die abstrakte Struktur der ZS transferiert wird. Um die Übersetzung zu bekommen, wird im Generierungsschritt die erzeugte abstrakte Repräsentation von Strukturen in eine natürlichsprachliche Ausgabe von Sätzen in der ZS umgewandelt, wie es in Abbildung 6 veranschaulicht wird (vgl. Werthmann & Witt 2014: 89f).

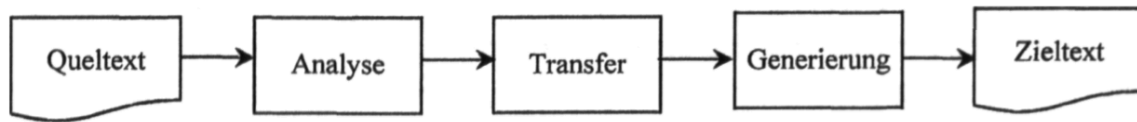


Abbildung 6: Schema für Transferbasierte Übersetzung (Carstensen et al. 2010: 646)

Der eindeutige Vorteil des Transfer-Ansatzes gegenüber der direkten Übersetzung besteht darin, dass im Übersetzungsprozess sowohl eine syntaktische als auch eine semantische Analyse stattfindet, wodurch deutlich bessere Ergebnisse erzielt werden. Es gilt festzuhalten, dass auch dieser Ansatz nicht in der Lage ist, „perfekte“ Übersetzungen zu liefern. Wie die Praxis zeigt, gibt es einen Punkt, ab dem die höhere Komplexität des Prozesses nicht zur Verbesserung der Übersetzungsqualität beiträgt, sondern oft sogar zu neuen Fehlern führen kann (vgl. Stein 2009: 8; Werthmann & Witt 2014: 89f).

3.2.1.3 Interlingua-Ansatz

Der dritte regelbasierte Ansatz ist die Interlingua-Übersetzung, bei welcher der Aufwand für die Analyse der AS im Vergleich zu den anderen zwei Ansätzen am größten ist. Der Übersetzungsprozess besteht im Wesentlichen aus zwei Schritten: Analyse und Generierung. Ähnlich wie beim Transfer-Ansatz erfolgt zuerst die Segmentierung und die Analyse der AS. Danach wird versucht, eine sprachenunabhängige Zwischenrepräsentation des ATs zu erzeugen, die auch als Interlingua bezeichnet wird. Diese Kodierung wird dann im Generierungs-Schritt in die ZS umgewandelt, um eine natürlichsprachliche Übersetzung zu bekommen, wie es in Abbildung 7 dargestellt wird (vgl. Werthmann & Witt 2014: 90).

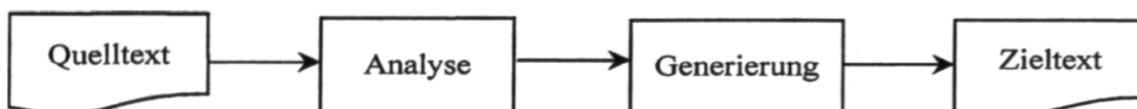


Abbildung 7: Schema für Interlingua-basierte Übersetzung (Carstensen et al. 2010: 646)

Dieser Ansatz wird oft als utopisch bezeichnet, da er auf der Annahme beruht, es wäre möglich, alle sprachlichen Informationen auf eine völlig sprachenunabhängige Art zu kodieren. Dazu müssten alle Unschärfen und Mehrdeutigkeiten einer Sprache aufgelöst werden, um sich vollständig von der Struktur der AS lösen zu können und um einen neuen, vom Ausgangstext unabhängigen, aber gleichwertigen Text in der ZS generieren zu können. Leider wurde bis heute noch keine ideale Repräsentationssprache entwickelt, die einerseits über eigene lexikalische und strukturelle Elemente verfügen würde und andererseits alle Ausdrucksmöglichkeiten von allen Sprachenpaaren berücksichtigen und formalisieren könnte. Die Entwicklung einer solchen abstrakten universalsprachlichen Repräsentation könnte die Arbeit sämtlicher Übersetzungssysteme wesentlich erleichtern. Heute gibt es diesen Ansatz aber lediglich in experimentellen Systemen, die nur für bestimmte Domänen oder für Zielsprachen mit kontrolliertem Vokabular und stark vereinfachten Satzstrukturen konzipiert sind (vgl. Werthmann & Witt 2014: 90f).

3.2.2 Statistikbasierte Ansätze

Ein wesentlicher Nachteil der regelbasierten Ansätze besteht darin, dass ihre Realisierung mit einem großen Aufwand verbunden ist. Sie basieren nämlich auf zweisprachigen Wörterbüchern und Regelsammlungen, die zuerst manuell erstellt und zusammengestellt werden müssen. Aus diesem Grund gewannen seit 1989 statistikbasierte maschinelle Übersetzungsstrategien (engl. Statistical Machine Translation, SMT) – mit denen eine größere Genauigkeit mit geringem Aufwand erreicht werden kann – immer mehr an Bedeutung (vgl. Tripathi & Sarkhel 2010). Statistische Ansätze arbeiten mit parallelen und alignierten Textkorpora, aus denen die Informationen mit Hilfe von empirischen Auswertungsmethoden extrahiert werden. Es handelt sich dabei grundsätzlich um zwei- oder mehrsprachige Korpora, die über den gleichen Inhalt verfügen. Im Laufe der sogenannten Alignierung werden inhaltsidentische Entsprechungen zwischen Korpora auf Wort-, Phrasen- oder Satzebene gesucht und markiert. Diese Markierungen können sowohl manuell als auch automatisch oder halbautomatisch erstellt werden. Die manuelle Alignierung kann zwar die Qualität der Übersetzung deutlich verbessern, ist aber umständlich und nicht effizient, wenn es sich um eine große Menge von Texten handelt. Aus Sicht der modernen MD-Systeme ist die automatische Alignierung relevant, bei der das Ausgangsdokument und seine Übersetzung in Übersetzungseinheiten segmentiert werden, die mit Hilfe linguistischer und statistischer Algorithmen einander zugeordnet und zusammengefügt werden. Obwohl dieses Verfahren sehr effizient ist, können Fehler und Überlappungen zwischen den Übersetzungseinheiten vorkommen. Da es oft mehrere Übersetzungsmöglichkeiten für ein Wort bzw.

für einen Satz gibt, besteht die Möglichkeit, dass nicht die entsprechende und im Kontext passende Übersetzung gefunden wird. Außerdem kann sich sogar die Reihenfolge der Textfragmente bei der Übersetzung ändern oder verschieben (vgl. Carstensen et al. 2010; Werthmann & Witt 2014: 91f).

Grundsätzlich ist die Funktionsweise aller korpusbasierten MÜ-Ansätze ähnlich: Es wird zuerst nach entsprechenden Beispielen aus dem Korpus der AS gesucht, danach werden die am besten übereinstimmenden zielsprachlichen Äquivalente geliefert. Korpusbasierte Übersetzungsstrategien lassen sich in weitere Ansätze unterteilen: Es wird beispielsweise zwischen reinstatistik-, beispiel- und kontextbasierten MÜ-Strategien unterschieden (vgl. Ramlow 2009; Stein 2009; Werthmann & Witt 2014: 92).

3.2.2.1 Reinstatistikbasierte Übersetzung

Der rein statistisch arbeitende MÜ-Ansatz ist eines der meist verbreiteten korpusbasierten Lernverfahren. Das System stützt sich auf die Auswertung von großen parallelen und alignierten Sprachdaten und berechnet die Wahrscheinlichkeit der verschiedenen Übersetzungen entweder auf Wort- oder auf Phrasenebene. Hierbei werden Aspekte wie Worthäufigkeit, Position des Wortes im Satz, Satzlänge usw. berücksichtigt und ein Übersetzungsmodell erstellt, „das eine Liste von Übersetzungsmöglichkeiten für jeden ausgangssprachlichen Term mit der berechneten Wahrscheinlichkeit seiner zielsprachlichen Entsprechungen enthält.“ (Werthmann & Witt 2014: 92f) Dieser Ansatz basiert ausschließlich auf den statistischen Wahrscheinlichkeitsberechnungen, in den Prozess werden weder Wörterbücher noch linguistische Methoden einbezogen. Abhängig davon, auf welcher Ebene die Erstellung des Übersetzungsmodells erfolgt, wird zwischen wort- und phrasenbasiertem Ansatz unterschieden (vgl. Werthmann & Witt 2014: 92f).

3.2.2.2 Wortbasierte statistische MÜ

Die Grundidee des Modells der wortbasierten statistischen Übersetzungsstrategie geht auf das Projekt Candide von IBM zurück, das in den 1980er- und 1990-er Jahren durchgeführt wurde und den Satz nicht als Ganzes, sondern geteilt, auf Wortebene untersucht (vgl. Koehn 2010: 81). Wie das Übersetzungsmodell auf Wortbasis aussieht, wird in Tabelle 2 am Beispiel des deutschen Satzes „*Das Haus ist klein*“ veranschaulicht. Für jedes deutsche Wort werden fünf Übersetzungsmöglichkeiten ins Englische angegeben, denen jeweils eine Wahrscheinlichkeit zugeordnet wird. Der relative Wert der Wahrscheinlichkeit (zwischen 0 und 1) lässt sich

errechnen, indem die Zahl der einzelnen Möglichkeiten durch die Gesamtzahl aller Möglichkeiten geteilt wird. Es zeigt sich, dass die wahrscheinlichste Übersetzung für das deutsche Wort „Haus“ das englische „house“ ist: In 80% der Fälle wurde das erste durch das zweite ersetzt. Dieser Ansatz geht davon aus, dass einem Wort der AS immer ein bestimmtes Wort der ZS entspricht, was in der Wirklichkeit aber nicht immer der Fall ist. Bereits das einfache Beispiel in Tabelle 2 zeigt, dass es für einzelne Wörter mehrere Übersetzungsmöglichkeiten gibt (vgl. Werthmann & Witt 2014: 93f).

Tabelle 2: Berechnung der Übersetzungswahrscheinlichkeiten für jedes Wort im Satz „Das Haus ist klein“ nach Koehn (2010: 84)

das		Haus		ist		klein	
Übersetzung (Ü)	Wahrscheinlichkeit (P)	Ü	P	Ü	P	Ü	P
the	0.7	house	0.8	is	0.8	small	0.4
that	0.15	building	0.16	's	0.16	little	0.4
which	0.075	home	0.02	exists	0.02	short	0.1
who	0.05	household	0.015	has	0.015	minor	0.06
this	0.025	shell	0.005	are	0.005	petty	0.04

Ein großer Nachteil der wortbasierten Übersetzungsstrategie besteht darin, dass sie auf der Annahme beruht, dass die einzelnen Wortübersetzungen voneinander unabhängig existieren. Sogar die Kontextinformationen werden bei der Wahrscheinlichkeitsberechnung außer Acht gelassen, obwohl sie die Übersetzung der einzelnen Wörter stark beeinflussen können. Wenn nicht Wörter, sondern Phrasen – die aus mehreren Wörtern bestehen – als Übersetzungseinheiten betrachtet werden, können deutlich bessere Ergebnisse erzielt werden (vgl. Werthmann & Witt 2014: 94).

3.2.2.3 Phrasenbasierte statistische MÜ

Wie schon vorher erläutert, werden bei der phrasenbasierten MÜ statt einzelnen Wörtern Phrasen, also Mehrwortsequenzen, als Übersetzungseinheiten angesehen. Diese Segmentierung erfolgt in erster Linie auf linguistischer und nicht unbedingt auf syntaktischer Ebene (vgl. Koehn 2010: 128). Nachdem die einzelnen Phrasen aus der AS in die ZS übersetzt wurden, stehen sie weiterhin in derselben Reihenfolge wie in der AS. Es darf aber nicht vergessen werden, dass die Wortreihenfolge in der AS und ZS stark voneinander abweichen kann, deswegen werden

die zielsprachigen Übersetzungseinheiten mit Hilfe eines Sprachmodells in die grammatikalisch richtige Reihenfolge geordnet. Dieses Modell wird auf Basis von einsprachigen Korpora gebildet und hat die Aufgabe, aus den übersetzten Phrasen die wohlgeformte Wortfolge der ZS zu berechnen. Ein bedeutender Vorteil der phrasenbasierten gegenüber der wortbasierten Übersetzungsstrategie ist, dass hier ein einzelnes Wort mit mehreren Wörtern übersetzt werden kann und umgekehrt. Da auch der Kontext in den Übersetzungsprozess einbezogen wird, können Mehrdeutigkeiten besser behandelt und gelöst werden. (vgl. Carstensen et al. 2010; Werthmann & Witt 2014: 94).

Wie auch in Abbildung 8 dargestellt wird, besteht der Übersetzungsprozess reinstatistikbasierter Ansätze grundsätzlich aus zwei Hauptschritten: Im ersten Schritt erfolgt die Decodierung des Ausgangstextes mit Hilfe des Übersetzungsmodells, das eine Liste mit den unterschiedlichen Übersetzungsmöglichkeiten und ihrer Wahrscheinlichkeit ermittelt. Im zweiten Schritt werden die kontextuell wahrscheinlichsten Übersetzungsalternativen ausgewählt und in eine grammatikalisch möglichst korrekte zielsprachige Wortreihenfolge geordnet (vgl. Werthmann & Witt 2014: 94f).



Abbildung 8: Schema für die statistische MÜ (Carstensen et al. 2010: 650)

Im Vergleich zu regelbasierten Ansätzen kann man feststellen, dass reinstatistische Ansätze sowohl im Kosten- als auch im Zeitaufwand günstiger sind, unter anderem, weil sie auf die mühsame manuelle Erstellung von komplizierten Grammatikregeln und auf die Auflistung ihrer Ausnahmen verzichten. Da die meisten Algorithmen in der statistischen MÜ sprachunabhängig sind, ist die Erweiterung solcher Übersetzungssysteme um weitere Sprachen relativ einfach: Durch Hinzunahme von neuen, alignierten bzw. parallelen Textdaten für das gegebene Sprachenpaar wird das System automatisch um neue Sprachen erweitert. Statistikbasierte MÜ ist mit einer schnellen Realisierung, mit geringen Investitionskosten und mit hoher Funktionssicherheit verbunden und ist daher sehr effizient. Dank der vielen Vorteile wurden die bekanntesten Online-Übersetzungssysteme der Firmen *Google* und *Microsoft* bis 2016 auf einem statistischen Ansatz aufgebaut (vgl. Werthmann & Witt 2014: 95f; McGuire 2019).

Natürlich darf man die Nachteile des reinstatistischen Ansatzes auch nicht unerwähnt lassen. Die Erstellung und Aufbereitung von großen Mengen von zuverlässigen Textkorpora, aus denen Wahrscheinlichkeitsberechnungen erhalten werden können, erfordert einen großen Aufwand an Zeit und Energie. Nach Bennet & Gerber (2003) ist ein rund eine Million Wörter umfassender Korpus von zweisprachigen Satzpaaren nötig, um ein MÜ-System universell einsetzen zu können (vgl. Werthmann & Witt 2014: 96).

3.2.2.4 Beispielbasierte MÜ

Bei der beispielbasierten MÜ (engl. Example Based Machine Translation, EBMT) werden Segmente aus bereits übersetzten Texten abgerufen und wiederverwendet. Es wird mit bilingualen parallelen Textdaten gearbeitet, die von Fachleuten übersetzt wurden und nicht von solchen, die lediglich über Sprachkenntnisse verfügen. Auf diese Weise kann man davon ausgehen, dass die Verknüpfung von AT und ZT auf Basis von Fachwissen erfolgte bzw. dass im Laufe des Übersetzungsprozesses bewusste translatorische Entscheidungen getroffen wurden und deshalb die Übersetzung zuverlässig genug ist, um sie wiederzuverwenden zu können. Im Laufe des beispielbasierten maschinellen Übersetzungsprozesses wird nach Segmenten gesucht, die dem zu übersetzenden Satz oder Satzteil ähneln oder mit ihm identisch sind (vgl. Werthmann & Witt 2014: 96).

Die Idee, die von Fachleuten bereits übersetzten und gespeicherten Texte im maschinellen Übersetzungsprozess wiederzuverwenden, wurde erstmals von Makoto Nagao im Jahr 1981 vorgeschlagen, sie wurde aber erst drei Jahre später veröffentlicht (vgl. Somers 2003). Durch Einbeziehung von menschlichen ÜbersetzerInnen kann man die Qualität der Übersetzung steigern und das in den bilingualen Textdaten kodierte Fachwissen extrahieren und wieder benutzen (vgl. Werthmann & Witt 2014: 96).

Der Übersetzungsprozess besteht aus vier Schritten: Im ersten Schritt werden die für die Übersetzung relevanten Beispiele mit Hilfe automatischer Textalignierung aus einem bilingualen Korpus extrahiert und ermittelt. Diese ausgewählten Texteinheiten werden danach in einer Datenbank gespeichert. Im zweiten Schritt werden die Eingaben in der Datenbank verwaltet bzw. für die Ausgabe und den Abruf von Beispielen gesorgt. Im dritten Schritt werden die Einträge, die mit der Eingabe am ähnlichsten sind, ausgewählt und für die automatische Übersetzung angewendet. Im letzten Schritt werden die aus der Datenbank erstellten Teilübersetzungen zu einer korrekten Übersetzung rekombiniert (vgl. Werthmann & Witt 2014: 96f).

Ein wesentlicher Vorteil dieses Ansatzes besteht darin, dass – im Unterschied zu reinstatistischen MÜ-Ansätzen – auch das in Beispielen eingeschlossene Sprachwissen in den

automatischen Übersetzungsprozess einbezogen wird. Es wird auf bereits vorhandene und gespeicherte Übersetzungen (von Sätzen, Redewendungen bzw. Sequenzen von Wörtern) zurückgegriffen und nach möglichst genauen Übereinstimmungen gesucht, die danach zu vollständigen Texten in der ZS zusammengesetzt werden. Werden keine Übereinstimmungen in den Korpusdaten gefunden, wird auch keine Übersetzung geliefert, woraus sich der größte Nachteil der beispielbasierten Übersetzungsstrategie ergibt (vgl. Werthmann & Witt 2014: 97).

Da man bei diesem Ansatz keine zeitaufwändigen sprachspezifischen Regeln angeben muss, ist der Aufbau eines beispielbasierten MÜ-Systems mit parallelen Korpora und Lexika relativ schnell zu bewerkstelligen und mit geringem Aufwand verbunden (vgl. Werthmann & Witt 2014: 97).

3.2.2.5 Kontextbasierte MÜ

Ein bedeutender Nachteil sowohl statistikbasierter als auch beispielbasierter Ansätze besteht darin, dass für ihre Umsetzung eine große Menge von parallelen Textdaten benötigt wird, von denen die Qualität und Zuverlässigkeit der Übersetzung abhängt. Natürlich ist die Verfügbarkeit solcher Korpora stark sprachabhängig, da es Sprachenpaare gibt, die öfter übersetzt werden und für die daher auch viele parallele Texte zur Verfügung stehen; jedoch gibt es Sprachenpaare, zu denen nur wenige oder gar keine Übersetzungen vorliegen (vgl. Werthmann & Witt 2014: 98).

Beim kontextbasierten Ansatz (engl. Context Based Machine Translation, CBMT) liegt dieses Problem nicht vor; für seine Umsetzung sind keine parallelen Korpora erforderlich. Die kontextbasierte MÜ stützt sich auf umfangreiche einsprachige Korpora der ZS und auf zweisprachige Vollformellexika. Zur Erstellung des monolingualen Textkorpus können einfach aus dem Internet heruntergeladene Texte verwendet werden (vgl. Werthmann & Witt 2014: 98).

Im Übersetzungsprozess wird der AT zuerst in sogenannte N-Gramme (Einheiten, die zwischen 4 und 8 Wörter lang sind) zerteilt. Danach werden mit Hilfe des Wörterbuches für jedes Wort alle möglichen Übersetzungsvarianten ermittelt, übersetzt und in zielsprachliche N-Gramme übertragen. Aus diesen Varianten wird für jedes Wort diejenige ausgewählt, die größere oder längere Treffer im Korpus hat; der Rest wird verworfen. Werden vom System keine überschneidenden N-Gramme im Korpus gefunden, kann das System mit einem Synonym-Generator nach weiteren möglichen Übersetzungen suchen (vgl. Carbonell et al. 2006). Da die Wörter in diesem Ansatz in ihrem Kontext übersetzt werden, können viele Ambiguitäten im Laufe des Übersetzungsprozesses gelöst werden (vgl. Werthmann & Witt 2014: 98).

3.2.3 Neuronale MÜ

Die in den vergangenen Jahrzehnten erreichten Fortschritte und Entwicklungen im Bereich der künstlichen Intelligenz haben in den 2010er Jahren einen Paradigmenwechsel in der maschinellen Übersetzungsbranche herbeigeführt: Es entstand eine neue und moderne maschinelle Übersetzungsstrategie, nämlich der neuronale MÜ-Ansatz (engl. Neural Machine Translation, NMT). NMT-Systeme waren ursprünglich für die Unterstützung statistikbasierter Übersetzungsansätze gedacht, haben diese jedoch in kurzer Zeit gänzlich ersetzt. Während bei einer Konferenz über MÜ im Jahr 2015 nur ein einziges Übersetzungssystem vorgestellt wurde, das gänzlich mit dem neuronalen Ansatz arbeitete, waren es bei einer ähnlichen Veranstaltung im Jahr 2017 (Conference on Machine Translation, Kopenhagen) fast ausschließlich NMT-Systeme, über die referiert wurde (vgl. Koehn 2017: 6f).

Der Ansatz der neuronalen MÜ basiert auf künstlichen Netzen, die dem neuronalen Aufbau des menschlichen Gehirns nachempfunden wurden. Künstliche neuronale Netze sind in der Lage, selbstständig aus der seit Jahren exponentiell zunehmenden Datenmenge (Big Data) zusätzliches Wissen zu extrahieren und dieses im Laufe des Übersetzungsprozesses zu verwenden. Im Wesentlichen handelt es sich hier um Repräsentationslernen, wobei „jede Schicht des neuronalen Netzes eine Repräsentation von der vorherigen Schicht lernt.“ (Patel 2020: 174) Von Schicht zu Schicht kommen immer detailliertere und nuanciertere Darstellungen vor, wobei auch die grammatikalischen Zusammenhänge zwischen Ausgangs- und Zielsprache implizit erfasst werden (vgl. Lionbridge 2017; Patel 2020: 174).

Bei neuronalen Netzen wird zwischen zwei Formen unterschieden: *flache Netze* haben weniger Schichten, während *tiefe Netze* mit vielen Schichten arbeiten. Heutzutage ist der Begriff *Deep Learning* (tiefes Lernen) in aller Munde: Diese Art des maschinellen Lernens bezieht sich auf die tiefen (vielschichtigen) neuronalen Netze, mit denen es arbeitet. Während *flache Netze* nicht besonders effektiv sind, ist *Deep Learning* ein unglaublich leistungsfähiges Verfahren und zählt zu den meist beforschten Bereichen im maschinellen Lernen (vgl. Patel 2020: 174).

Grundsätzlich kann man sich die Funktionsweise neuronaler Netze folgenderweise vorstellen: Anders als bei der statistikbasierten MÜ, wird bei NMT-Systemen immer ein vollständiger Satz betrachtet. Mathematisch gesehen erfolgt die Übersetzung – im Gegensatz zu SMT-Systemen – nicht auf einer linearen Art und Weise; dies bedeutet, dass der Output nicht direkt auf Basis des Inputs generiert wird (vgl. Schmalz 2019: 199). Zwischen der *Input-Schicht* und der *Output-Schicht* befinden sich sogenannte verdeckte Schichten (*hidden layers*), die als

unsichtbare Zwischenberechnungen aufgefasst werden können (Abbildung 9). Jede Schicht besteht aus Knoten, die auch Neuronen oder Einheiten genannt werden. Die Neuronen jeder Schicht werden mittels Informationsvektoren mit den Neuronen der nächsten Schicht verbunden (vgl. Patel 2020:174). Diese Vektoren verknüpfen also die Informationen zu jedem Wort mit Hilfe der benachbarten Wörter. Manchmal können das Hunderte von Informationen pro Wort sein, wodurch eine sehr hohe Genauigkeit erreicht werden kann (vgl. Lionbridge 2018).

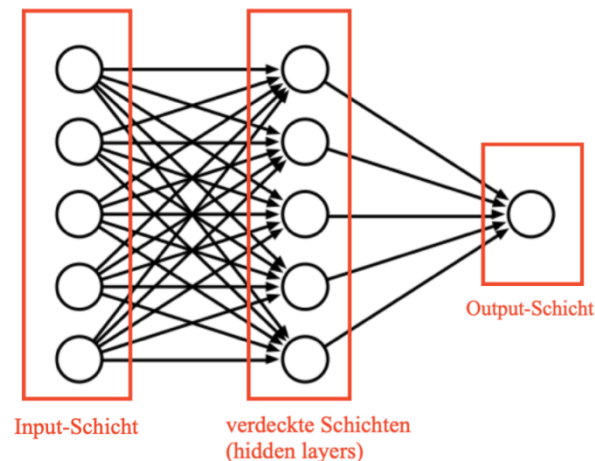


Abbildung 9: Neuronales Netz mit verdeckten Schichten (hidden layers) (Laki 2018: 170, übersetzt aus dem Ungarischen)

Um die für die Übersetzung relevanten Informationen herausfiltern zu können, wird vom System ein sogenanntes Attention-Modell verwendet (vgl. Lionbridge 2018). Die Übersetzungen, die mit NMT-Systemen erzeugt wurden, klingen deutlich natürlicher und flüssiger als die Übersetzungen von SMT-Systemen. Auch die Satzstruktur ist wesentlich besser: Sogar bei Sprachenpaaren, bei denen die Wortreihenfolge stark unterschiedlich ist, werden weniger Syntaxfehler registriert (vgl. Schmalz 2019: 199).

In den letzten Jahren hat Google mit einem noch aktuelleren Ansatz, basierend auf Recurrent Neural Networks (RNN), gearbeitet. Diese Netze funktionieren im Grunde genommen genauso wie andere neuronale Netze, sind aber – wie auch der Name verrät – rückgekoppelt: Die Knoten jeder Schicht sind auch mit den Knoten ihrer eigenen oder der vorhergegangenen Schicht verbunden, nicht nur mit Neuronen der nächsten Schicht (vgl. Breitenbach 2018).

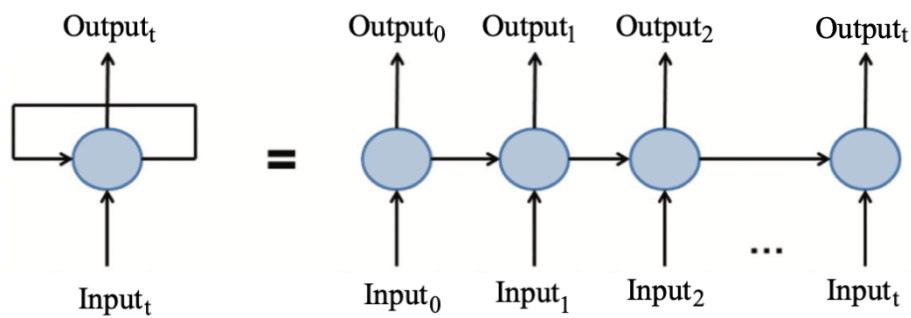


Abbildung 10: Schematische Abbildung einer Schicht des rekurrenten neuronalen Netzes (Laki 2018: 171, übersetzt aus dem Ungarischen)

An dieser Stelle soll erwähnt werden, dass trotz der vielen Vorteile der NMT-Systeme statistikbasierte Ansätze nicht komplett von solchen, die auf neuronalen Netzwerken beruhen, verdrängt wurden. Im Verlauf des Jahres 2017 integrierten fast alle großen Anbieter von Übersetzungswerkzeugen (wie Google, Microsoft/Bing, Systran, SDL, Yandex usw.) die NMT-Technologie in ihre Systeme, der SMT-Algorithmus wurde dadurch jedoch nicht automatisch ersetzt: Es wurden hybride Systeme entwickelt, in denen die Vorteile beider Ansätze miteinander vereint werden (vgl. Schmalz 2019: 199).

3.2.4 Probleme der MÜ

Dass die MÜ-Systeme trotz der Vielzahl von Methoden und Ansätzen nicht die Leistung der HumanübersetzerInnen erzielen können, liegt nicht in erster Linie an den angewendeten Technologien, sondern vielmehr an der Komplexität der natürlichen Sprache und ihrer Verarbeitung. Im Folgenden werden einige Herausforderungen der MT-Systeme vorgestellt, die den Übersetzungsprozess stark erschweren können (vgl. Werthmann & Witt 2014: 83).

Eines der zentralen und wichtigsten Probleme der MÜ, das auf verschiedenen Spracherebenen auftreten kann, ist die Ambiguität, also die Mehrdeutigkeit der Sprache. Dieses Phänomen ergibt sich daraus, dass viele Wörter in verschiedenen Kontexten unterschiedliche Bedeutungen haben können.

Es werden grundsätzlich zwei Arten von Ambiguität voneinander unterschieden: die lexikalische und die strukturelle. Erstere bezieht sich auf Wortformen, die für zwei oder mehrere Begriffe stehen. Solche mehrdeutigen Wörter werden auch als Homonyme bezeichnet. (vgl. Werthmann & Witt 2014: 83).

z.B.

dt. *die Bank* – engl. *bench*

dt. *die Bank* – engl. *bank*

Von struktureller Mehrdeutigkeit spricht man, „wenn ein Element nach dem Regelwerk einer Ausgangssprache von verschiedenen Elementen regiert werden kann und dies dem Syntaxmuster der Zielsprache nicht entspricht.“ (Werthmann & Witt 2014: 83) Der Satzbau ermöglicht also mehrere unterschiedliche Interpretationen desselben Satzes, aus denen MÜ-Systeme nicht in jedem Fall die „korrekte“ Variante identifizieren können.

z.B.

dt. *Alte Männer und Frauen* blieben zurück. – it. *Uomini anziani e le donne sono rimasti indietro.*

dt. *Alte Männer und Frauen* blieben zurück. – it. *Uomini e donne anziani sono rimasti indietro.*

Im deutschen Beispielsatz ist es unklar, ob das Adjektiv „alte“ sich nur auf die Männer oder auch auf die Frauen bezieht. Wird er in eine Sprache übersetzt, in der – im Gegensatz zur deutschen Sprache, wo das Adjektiv immer vor dem Substantiv steht – das Adjektiv in den meisten Fällen hinter das Nomen gestellt wird, wie im Italienischen (und in vielen anderen romanischen Sprachen), kann dies zu Übersetzungsproblemen führen. In dem in dieser Arbeit untersuchten Sprachenpaar (Deutsch – Ungarisch) besteht dieses Problem nicht, da das Adjektiv in beiden Sprachen immer vor dem Substantiv steht. Diese Problemquelle kann aber bei der Untersuchung anderer Sprachenpaare relevant sein.

Die Übersetzung zwischen Sprachen, die große strukturelle Unterschiede zwischen ihren Grammatiken aufweisen, ist naturgemäß problematischer und komplexer als die Übersetzung von Sprachenpaaren mit ähnlichen Strukturen. Beispielsweise gelangen die mit Google Translate durchgeführten Übersetzungen zwischen dem Englischen und Spanischen wesentlich besser als solche zwischen dem Englischen und Japanischen (vgl. Schulz 2013). Was das Sprachenpaar Deutsch – Ungarisch angeht: Ungarisch ist, anders als die meisten europäischen Sprachen, nicht Teil der indoeuropäischen Sprachfamilie und unterscheidet sich stark in Bezug auf die grundlegenden Strukturen der deutschen Sprache. Diese Tatsache wird auch in der Fachliteratur, die von ExpertInnen des Bereiches verfasst wurde und die beiden Sprachen sowohl auf lexikalischer als auch auf grammatikalischer Ebene vergleicht, betont. Bassola (2008) beschäftigt sich in seinem Werk mit verschiedenen Ergebnissen aus deutsch-ungarischen kontrastiven Untersuchungen und erwähnt unter anderem wesentliche strukturelle Unterschiede zwischen

den Infinitivkonstruktionen, den Nominalkonstruktionen und der Wortstellung (vgl. Bassola 2008: 201f). An dieser Stelle soll auch angemerkt werden, dass – anders als in den flektierenden Sprachen – im Ungarischen die Bildung von Wortformen durch Agglutination erfolgt, also durch den Anhang der Suffixe an die Wortstämme in genau festgelegter Reihenfolge. Dank dieser agglutinierenden Wortbildung gibt es in der Sprachwissenschaft keinen Konsens bezüglich der Anzahl der Fälle. Manche reduzieren diese Zahl auf fünf, viel öfter wird aber von 15, 17 oder sogar 40 Fällen gesprochen². Im Deutschen gibt es dagegen nur vier Kasus (Nominativ, Genitiv, Dativ und Akkusativ). Die Anhäufung von Suffixen führt oft zu sehr langen Wörtern, die sehr viele Informationen beinhalten, die nur durch die Analyse des Wortstammes und seiner Endungen in ihrer richtigen Bedeutung erfasst werden können (vgl. ATS Sprachendienst 2004).

Außer der Komplexität einzelner Wörter können auch lange und komplizierte Sätze den Übersetzungsprozess erschweren. Da es sich in dem hier untersuchten Fall um spontan gesprochene Sprache handelt, können Wortwiederholungen, falsche Syntax oder grammatikalische Fehler im Ausgangstext auftreten, die ebenfalls eine Herausforderung für das MÜ-System bedeuten können.

3.3 Sprachsynthese

Jener Bestandteil, welcher MÜ-Systeme schließlich zu Dolmetsch-Systemen macht, also zu solchen, welche gesprochene Sprache einer AS in gesprochene Sprache einer ZS übertragen, ist die Sprachsynthese, in deren Rahmen die von der MÜ-Komponente generierten schriftlichen Texte in ein Sprachsignal umgesetzt werden. Dies kann mit dem Vorlesen eines Textes durch eine Person verglichen werden. Beim Vorlesen von Texten wird durch die Artikulation und Akzentuierung der Wörter klar, wie gut die Person die jeweilige Sprache beherrscht und ob sie überhaupt versteht, was sie vorliest. Werden Wörter falsch ausgesprochen, wird dies vor allem von MuttersprachlerInnen schnell bemerkt. Sprachsynthese-Systeme sind im Gegensatz zu Menschen zum heutigen Zeitpunkt noch nicht in der Lage, die in Lautsprache umgewandelten Texte zu „verstehen“; die Sprachsignalproduktion erfolgt lediglich durch Algorithmen (vgl. Pfister & Kaufmann 2017: 27, 194f).

Grundsätzlich kann zwischen zwei Formen der Sprachsynthese unterschieden werden: der TTS-Synthese (text-to-speech synthesis) und der CTS-Synthese (concept-to-speech synthesis). Bei der ersteren ist die Synthese neben der Sprachsignalproduktion auch für die richtige

² „Nominativ (alany eset), Akkusativ (tárgy eset), Dativ (részes eset), Possessiv (birtokos eset), Instrumental, Final, Transformativ, Superessiv, Sublativ, Delativ, Inessiv, Elativ, Illativ, Adessiv, Allativ stb. [= usw.] werden von den Ungarn als Fälle angesehen.“ (ATS Sprachendienst 2004)

Aussprache der Wörter und die angemessene Akzentstärke der Silben verantwortlich. Die CTS-Synthese hingegen wird in jenen Systemen eingesetzt, in welchen die Eingaben in einer hierarchisch strukturierten Form zu finden sind (z. B. wie bei einem Syntaxbaum). Diese Systeme verfügen bereits über linguistisches Wissen und generieren die akustischen Meldungen selbstständig, da der syntaktische Aufbau der Sätze bereits vor der Durchführung der Sprachsynthese bestimmt wurde (vgl. Pfister & Kaufmann 2017: 27).

Die Erzeugung des Sprachsignals erfolgt grundsätzlich in zwei Schritten. Zuerst wird der schriftlich vorliegende Text linguistisch verarbeitet und eine abstrakte Beschreibung des zu produzierenden Sprachsignals erstellt. Dies ist die sogenannte Transkriptionsstufe, die noch stimmunabhängig durchgeführt wird. In der zweiten, phonoakustischen Stufe wird die zuvor erstellte phonologische Darstellung – die die Intensität, die Lautdauer und die Grundfrequenz der einzelnen Laute bestimmt – in eine lautsprachliche, akustische Meldung umgewandelt. Dieser Teil, in dem die tatsächliche Sprachsignalproduktion erfolgt, ist bereits stimmabhängig (vgl. Pfister & Kaufmann 2017: 196, 199). In Abbildung 11 werden die zuvor dargelegten Schritte der Sprachsynthese veranschaulicht.

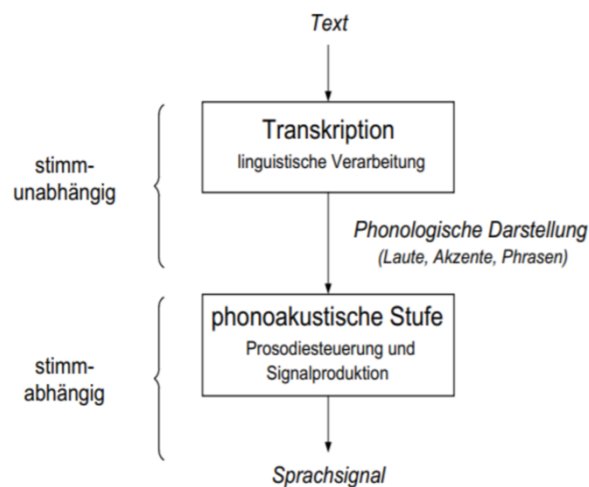


Abbildung 11: Ablauf der Sprachsynthese nach Pfister & Kaufmann (2017: 197)

Der Begriff der Sprachsynthese ist nicht mit dem der Sprachausgabe zu verwechseln; bei letzterer werden „lediglich“ Sprachsignale ausgegeben, die vorher von einer Person gesprochen und aufgezeichnet wurden. Im Gegensatz zur Sprachsynthese, mit der grundsätzlich jeder Text in Lautsprache umgesetzt werden kann, ist die Anzahl der möglichen akustischen Meldungen bei der Sprachausgabe immer auf die zuvor aufgenommenen Sprachsignale begrenzt (vgl. Pfister & Kaufmann 2017: 27).

Im folgenden Kapitel wird die von mir untersuchte Dolmetsch-Applikation *SayHi Translate* vorgestellt.

4. SayHi Translate

4.1 Überblick

Die App mit dem Slogan „*The Interpreter In Your Pocket*“ wurde 2009 von einem Start-Up-Unternehmen in Wethersfield (US) erfunden. Die Idee zur Entwicklung der App entstand im medizinischen Bereich: Nach Beobachtung des Krankenhausalltags wurde klar, dass viele PatientInnen wegen sprachlicher Schwierigkeiten nicht in der Lage waren, dem Fachpersonal ihre Beschwerden mitzuteilen. Dies führte dazu, dass viele, die der englischen Sprache nicht mächtig waren, keine entsprechende Behandlung im Krankenhaus erhielten (vgl. Hartford Business 2012). Deby Gould, Unternehmensleiterin und Mitbegründer von *SayHi Translate*, äußerte sich diesbezüglich folgendermaßen:

The idea for SayHi Translate was born from observation in a hospital for the medical necessity. There simply were too many patients unable to communicate with the medical staff in a way they could understand what was happening. There were too many patients leaving, unseen, after waiting hours for some type of interpretation services. (Hartford Business 2012)

Ursprünglich wurde das Programm „lediglich“ für die Erleichterung der Kommunikation zwischen medizinischem Personal und PatientInnen gedacht, nach einer dreijährigen Entwicklungsphase wurde es aber auch im iOS App Store der breiten Masse zur Verfügung gestellt. Anfangs konnte die App nur auf iPhone, iPad und iPod Touch kostenpflichtig heruntergeladen werden, heute ist sie aber kostenlos auch im Google Play Store und Amazon Appstore³ für alle iOS- und Android-Geräte verfügbar (vgl. SayHi Translate 2020)



Abbildung 12: Logo und App-Symbol von SayHi Translate (Screenshot)

³ An dieser Stelle soll erwähnt werden, dass die Applikation *SayHi Translate* von der Firma *Amazon* aufgekauft wurde und derzeit als *SayHi – an Amazon company* notiert wird.

Dank der ständigen Weiterentwicklung bietet *SayHi Translate* sprachübergreifende Kommunikation in derzeit 90 Sprachen und Dialekten an und gehört zu den meist verwendeten Dolmetsch-Applikationen für Smartphones (vgl. SayHi Translate 2020).

Sprachen	Dialekte
Afrikaans, Arabisch, Bosnisch, Bulgarisch, Dänisch, Deutsch, Englisch, Estnisch, Finnisch, Französisch, Griechisch, Haitianisch, Hebräisch, Hindi, Hmong, Indonesisch, Italienisch, Japanisch, Kantonesisch, Katalanisch, Koreanisch, Kroatisch, Lettisch, Litauisch, Malaiisch, Maltesisch, Mandarin (Taiwan), Mandarin (Vereinfachtes Chinesisch), Niederländisch, Norwegisch, Persisch, Polnisch, Portugiesisch, Rumänisch, Russisch, Serbisch, Slowakisch, Slowenisch, Spanisch, Suaheli, Schwedisch, Thailändisch, Tschechisch, Türkisch, Ungarisch, Ukrainisch, Urdu, Vietnamesisch, Walisisch	Arabisch: Algerien, Ägypten, Bahrain, Irak, Jemen, Jordanien, Katar, Kuwait, Libyen, Marokko, Oman, Palästina, Saudi-Arabien, Tunesien, VAE Englisch: Australien, Großbritannien, Indien, Kanada, Neuseeland, Südafrika, USA Französisch: Frankreich, Kanada Italienisch: Italien, Schweiz Mandarin: China, Taiwan Portugiesisch: Brasilien, Portugal Spanisch: Argentinien, Bolivien, Costa Rica, Dominikanische Republik, El Salvador, Guatemala, Honduras, Kolumbien, Mexiko, Nicaragua, Panama, Paraguay, Peru, Puerto Rico, Spanien, Uruguay, USA, Venezuela

Abbildung 13: Sprachen und Dialekte in SayHi Translate (Screenshot)

Da in medizinischen Settings oft persönliche, individuell identifizierbare Daten vorkommen, die sich entweder auf die PatientInnen selbst (Name, Adresse, Geburtsdatum, Sozialversicherungsnummer, Kartennummer usw.) oder auf ihren medizinischen Zustand beziehen, ist für *SayHi Translate* der Datenschutz besonders wichtig. Die App ist HIPAA-konform⁴; dies bedeutet, dass sie die PatientInnendaten als geschützte Gesundheitsdaten⁵ behandelt und mit ihnen nach sehr strengen Regelungen umgeht. SayHi arbeitet mit einer Datenredundanz, also werden die Daten mehrfach an verschiedenen Stellen gespeichert, damit sie im Falle eines Verlustes

⁴ „HIPAA steht für Health Insurance Portability and Accountability Act. Dieses US-Gesetz von 1996 regelt die Sicherheit und den Datenschutz im Zusammenhang mit geschützten Patientendaten (Protected Health Information, PHI) sowie den Zugriff der Patienten auf medizinische Datensätze. HIPAA wurde seit seiner Verabschiedung mehrfach angepasst, um neuen Technologien und den modernen Datenschutzrisiken Rechnung zu tragen.“ (Box for healthcare 2020)

⁵ „Geschützte Gesundheitsdaten sind individuell identifizierbare Daten, die sich auf den medizinischen oder psychischen Zustand eines Patienten, die Bereitstellung medizinischer Dienste oder deren Bezahlung (in der Vergangenheit, Gegenwart oder Zukunft) beziehen. PHI beinhalten ebenfalls allgemeine Identifikationsmerkmale wie Name, Adresse, Sozialversicherungsnummer und Geburtsdatum des Patienten.“ (Box for healthcare 2020)

wiederbeschafft werden können. Unter dem Punkt „SayHi Nutzungsbedingungen“ werden Datenschutzhinweise, Vertragsbedingungen, Richtlinien und Regeln zur Verwendung der App zusammengefasst. Für die Verwendung der App ist eine (stabile) Internetverbindung erforderlich, da sie offline nicht arbeitet.

Die App verwendet für das Dolmetschen verschiedene Vokabularverwaltungssysteme und Speicher wie z.B. Amazon Cloud oder Dragon Medical Direct⁶, es werden aber auch Fachvokabulare eingesetzt, die vom Unternehmen selbst entwickelt wurden.

4.2 User Interface

Die App ist transparent und logisch aufgebaut und verfügt über ein einfaches und modernes Design in den Farben Blau-Weiß. Sobald die App geöffnet wird, erscheint nach dem Logo der Startbildschirm, der grundsätzlich weiß und leer ist. Die zwei Sprachen werden im unteren Teil des Bildschirms mit einem blauen und grünen Mikrofon symbolisiert. Nach ein paar Sekunden wird auch die Anweisung „Loslegen – Wählen Sie unten ihre Sprachen aus und tippen Sie auf das Mikrofon“ sichtbar, wie in Abbildung 14 dargestellt wird.

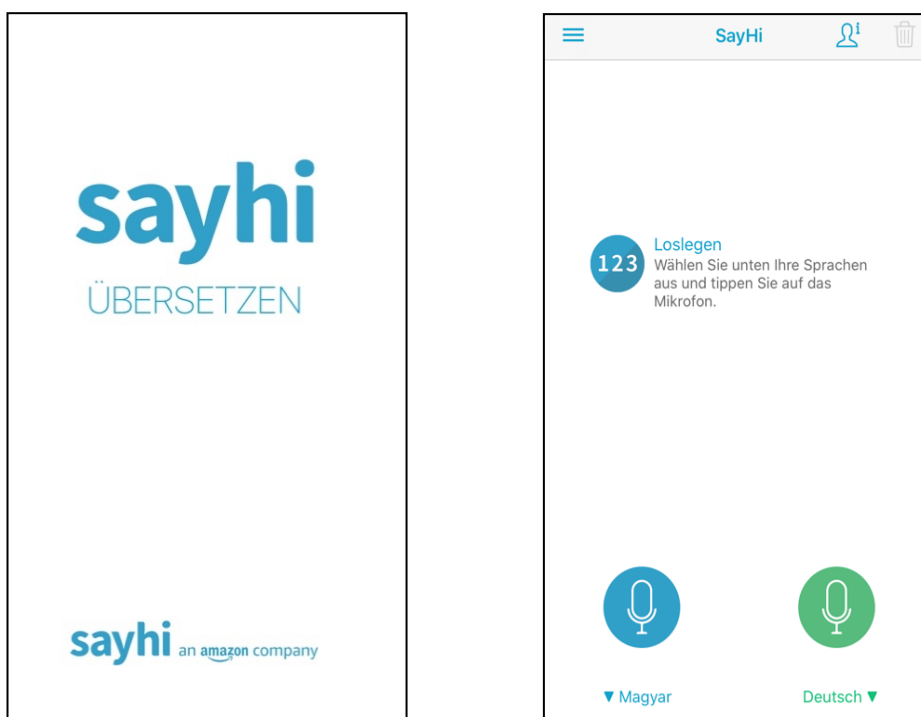


Abbildung 14: Startbildschirm von SayHi Translate (Screenshot)

⁶ Dragon Medical Direct ist ein Server für professionelle medizinische Spracherkennung, die über ein umfangreiches Fachvokabular verfügt und dank künstlicher Intelligenz die Erkennung von Fachbegriffen deutlich erleichtert (vgl. Voicepoint AG 2020).

Nach jeweils 2-3 Sekunden erscheint eine neue Anweisung auf dem Bildschirm, um die Funktionen der App schnell vorzustellen und so die Verwendung zu erleichtern. Dies ist vor allem bei der ersten Nutzung relevant.

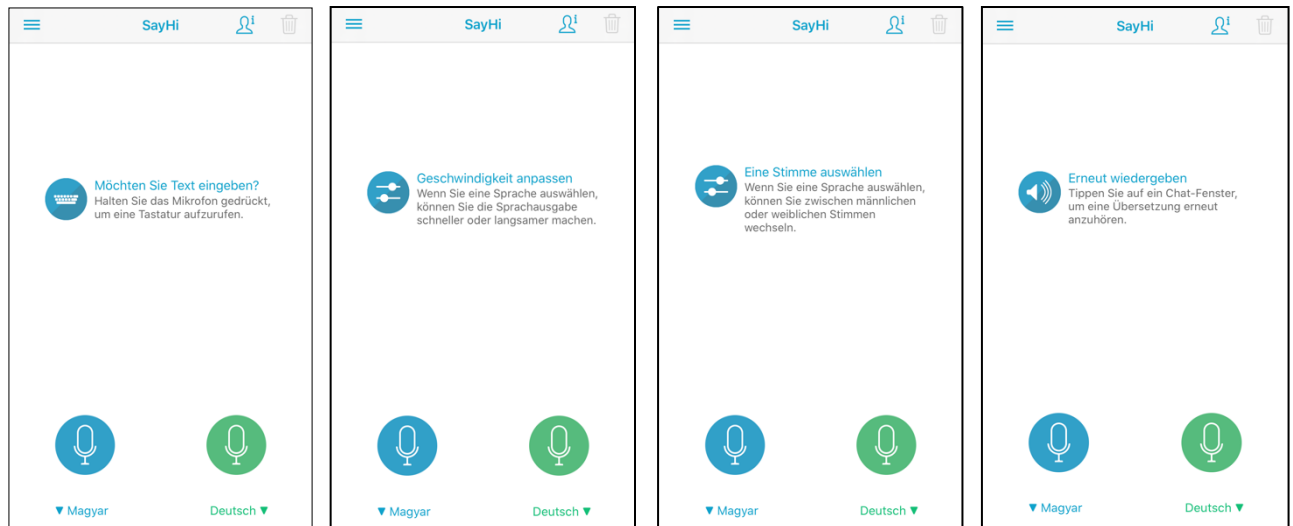


Abbildung 15: Anweisungen für die verschiedenen Funktionen von SayHi Translate (Screenshot)

Wie schon die Anweisungen verraten, gibt es die Möglichkeit, die Sprachausgabe an die eigenen Bedürfnisse anzupassen. Bei vielen Sprachen kann man die Geschwindigkeit der Sprachausgabe beschleunigen oder verlangsamen bzw. auswählen, ob eine männliche oder eine weibliche Stimme zu hören ist. Auch die Art und Weise der Spracheingabe kann hier eingestellt werden (Tippen oder Sprechen). Unter „Einstellungen“ können weitere Funktionen angepasst werden.

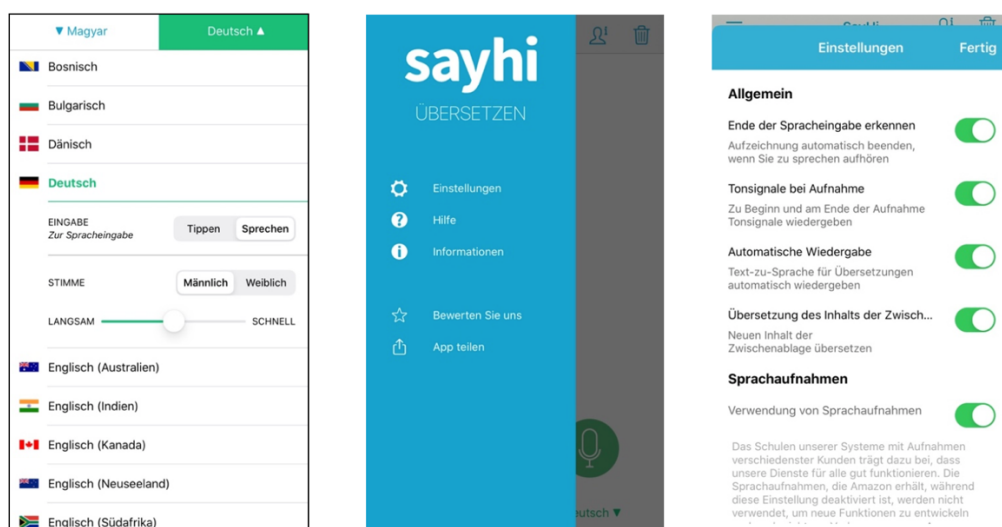


Abbildung 16: Einstellungen in SayHi Translate (Screenshot)

Nachdem die Sprachen ausgewählt und alle nötigen Einstellungen vorgenommen wurden, kann das sprachübergreifende Gespräch anfangen. Bevor man mit dem Sprechen beginnt, muss man das entsprechende Mikrofon-Symbol kurz drücken. Sobald die drei Punkte am Bildschirm erscheinen beginnt die automatische Spracherkennung. Die Aufzeichnung endet automatisch, sobald man mit dem Sprechen aufhört. Die mündlich dargebotene Informationseinheit wird danach in Schriftsprache umwandelt und von der MÜ-Komponente in die Zielsprache übersetzt. Die Übersetzung in der ZS erscheint zuerst in Textform und wird danach durch die Sprachsynthese auch als Sprachausgabe angeboten. Nach der maschinellen Übersetzung kann es vorkommen, dass *SayHi* um Feedback bittet, um sicherzustellen, dass das Gesagte von der App genau verstanden wurde. Dies dient auch dazu, dass sich die App an die Stimme „gewöhnt“ und sich ständig weiterentwickeln kann. In Abbildung 17 wird der Prozess der Spracherkennung und -verarbeitung dargestellt. Es wurde die ungarische Grußformel „Jó napot!“ mündlich eingegeben, in textueller Form erfasst, ins Deutsche übersetzt („Guten Tag“) und schließlich in Lautsprache ausgegeben. Nach dem Prozess wurde nachgefragt, ob der Ausgangstext von der App richtig verstanden wurde.

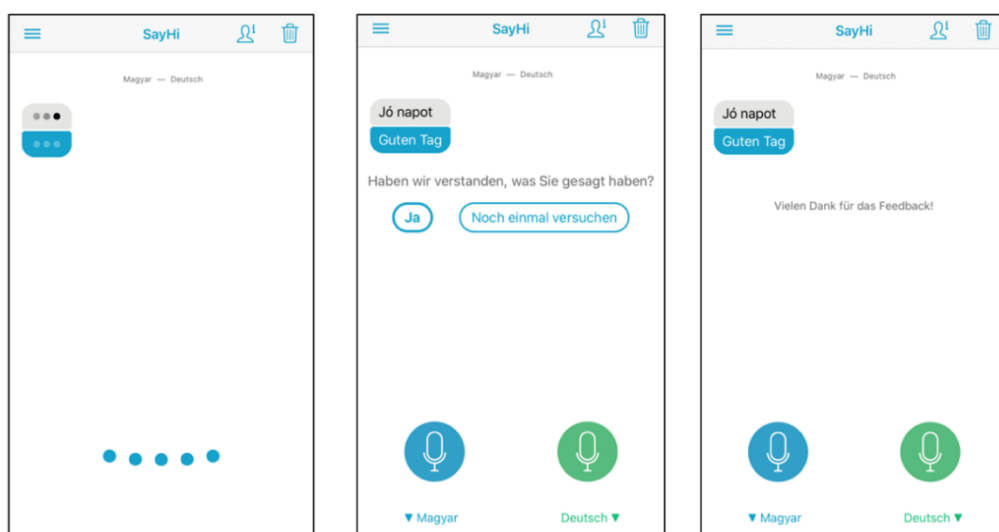


Abbildung 17: MD-Prozess (Screenshot)

Wie auch in Kapitel 3.1.1 erläutert wurde, kann es aufgrund verschiedener Störfaktoren, wie z.B. Umgebungsgeräuschen, Disfluenzen oder Dialekten, zu Problemen bei der Spracherkennung kommen. Außerdem kann es vorkommen, dass man zu früh zu sprechen beginnt und die Spracherkennung die ersten paar Wörter verpasst. In solchen Situationen ist die manuelle Korrekturfunktion der App sehr nützlich, da man auf diese Weise den erkannten Text mit Hilfe

einer Tastatur selbst korrigieren kann und die Wiederholung des Gesagten nicht nötig ist. Abbildung 18 zeigt, wie die manuelle Korrekturfunktion angewendet werden kann. Man hält den Text, den man korrigieren möchte, gedrückt und wählt danach die Option „Bearbeiten“ aus. Nach der Korrektur wird der AT wieder übersetzt und erneut in Lautsprache wiedergegeben.

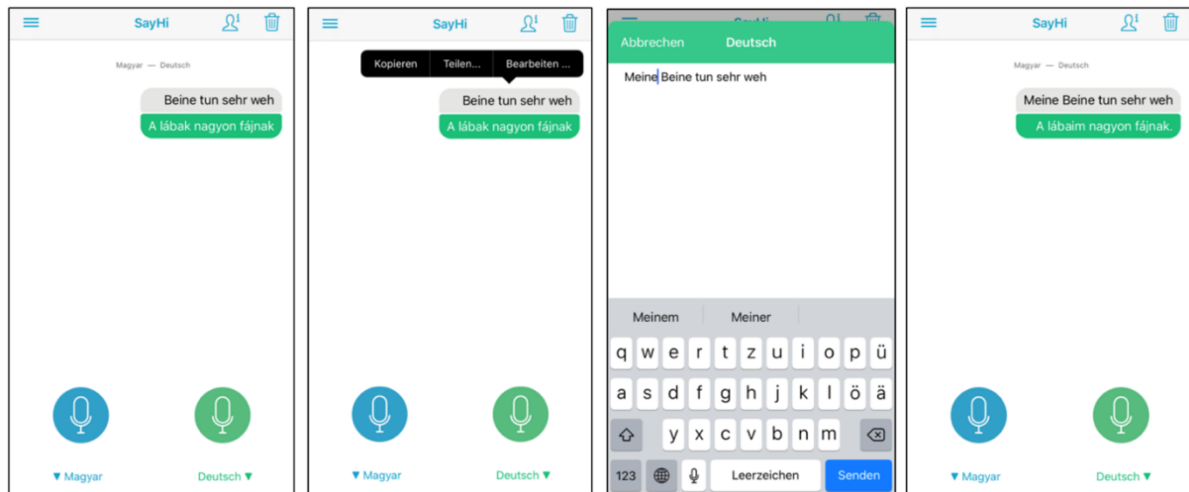


Abbildung 18: Manuelle Korrekturfunktion in SayHi Translate (Screenshot)

5. Dolmetschen im Gesundheitswesen

Die Überwindung von Sprach- und KulturbARRIEREN zwischen medizinischem Fachpersonal und PatientInnen mit Migrationshintergrund, die der deutschen Sprache nicht mächtig sind, stellt eine große Herausforderung im österreichischen Gesundheitsbereich dar. Immer häufiger werden mehrsprachige KlinikmitarbeiterInnen oder Verwandte der PatientInnen als Laien-DolmetscherInnen eingesetzt, jedoch qualifiziert Zweisprachigkeit allein nicht zum professionellen Dolmetschen. Diese spontanen und meist ungeplanten Dolmetscheinsätze sind zwar oft kostengünstig, können aber, wie bereits erwähnt, nicht nur gesundheitliche, sondern auch rechtliche Konsequenzen haben. Letztendlich geht es um Leib und Leben, eine fachlich korrekte und vollständige Wiedergabe der Inhalte ist also für die genaue Anamnese, Diagnose und Behandlung unerlässlich (vgl. Marquardt 2015). Im Folgenden wird die komplexe Aufgabe der DolmetscherInnen im medizinischen Bereich genauer untersucht und darauf eingegangen, welche Rollen DolmetscherInnen bei solchen Settings einnehmen können, um eine erfolgreiche Kommunikation zwischen medizinischem Personal und PatientInnen zu ermöglichen.

5.1 Die Rolle der DolmetscherInnen in medizinischen Settings

Die frühe Translationswissenschaft bemühte sich, eine passive, unsichtbare und unbeteiligte Rollenauffassung von Dolmetschenden als die gängige Praxis zu unterbreiten. Diese Vorstellung erwies sich aber vor allem in dialogischen Kontexten als falsch. An solchen translatorischen Interaktionen nehmen neben den primären Kommunikationsparteien (hier: medizinisches Personal und PatientInnen) DolmetscherInnen oft als eine dritte, aktive Partei teil und haben eine viel komplexere Aufgabe als das „bloße“ Transkodieren von Informationen und die Lieferung einer „wortwörtlichen“ Übersetzung. Alle Personen, die an der Kommunikation beteiligt sind, verfolgen ein bestimmtes Kommunikationsziel und nehmen daher eine entsprechende Rolle in der Gesprächssituation ein. Während sich PatientInnen ärztliches Verständnis und Gehör für ihre Beschwerden erwarten, möchten ÄrztInnen relevante Informationen für die Erstellung einer Diagnose bekommen. Dolmetschende hoffen, mit ihrer Leistung ein Verständnis zwischen den GesprächsteilnehmerInnen zu ermöglichen, und fungieren als ÜbermittlerInnen bzw. BrückenbauerInnen in der Kommunikation. Sie befinden sich dabei oft in höchst emotionalen Situationen, in denen Menschen Schicksalsschläge erleiden. Da aber neben den emotionalen auch zweckgebundene und fachliche Gesprächsinhalte vorkommen, sind Dolmetschende im Gesundheitswesen mit hohen Anforderungen konfrontiert. Während des Gesprächs

bedienen sich ÄrztInnen meistens semiprofessioneller Rede, mit welcher aber institutions- und prozessgebundenes Wissen vermittelt wird. Dieses Wissen wird von DolmetscherInnen für eine erfolgreiche Kommunikation vorausgesetzt (vgl. Angelelli 2004; Havelka 2018: 76ff; 80, Niska 2002).

Um eine erfolgreiche Kommunikation zwischen den GesprächspartnerInnen zu ermöglichen, müssen Dolmetschende über eine Reihe von Kompetenzen verfügen, die nach Angelelli (2004: 11) auf folgende Weise zusammengefasst werden können:

- communicating cultural gaps as well as linguistic barriers
- communicating affect nuances as well as the content of the message
- establishing trust between all parties to the conversation
- facilitating mutual respect
- putting the parties at ease during the conversation
- creating more balance (or imbalance) during the conversation (by aligning with one of the parties)
- advocating for or establishing alliances with either party
- managing the requested and given information

Man sieht also, dass die Aufgaben von DolmetscherInnen stark über die spezifischen Sprachkompetenzen hinausgehen. Sie müssen mit den unterschiedlichen kulturell bedingten Wertvorstellungen und Normen umgehen, die Machtverhältnisse gut einschätzen und die Interessen beider Gesprächsparteien vermitteln können. Sie erkennen und filtern die wichtigen Informationen und formulieren sie so, dass diese bei den RezipientInnen auch problemlos ankommen (vgl. Angelelli 2004: 11; Havelka 2018: 81f).

Holz-Mänttari (1984) befasste sich mit der Frage der Gesprächskoordination und stellte fest, dass DolmetscherInnen während der Interaktion eine aktive Rolle einnehmen und das Gespräch auf einer metakommunikativen Ebene bewusst koordinieren. Eine gute Koordination von SprecherInnenwechsel bzw. Sprechzeit gehört zu den wichtigsten Aufgaben der DolmetscherInnen in dialogischen Settings und ist gleichzeitig ein wesentlicher Faktor für eine gelungene Kommunikation. Zusätzlich steuern DolmetscherInnen auch den Themenwechsel und greifen sogar ein, um Missverständnisse zu vermeiden oder Hintergrundinformationen zu liefern. DolmetscherInnen müssen außerdem in der Lage sein, mit kommunikativen Störfaktoren wie z.B. Stille, Pausen oder gleichzeitigem Sprechen bewusst umzugehen (vgl. Havelka 2018: 85f, Holz-Mänttari 1984: 53).

Leanza (2007) erweitert das Rollenverständnis von Dolmetschenden im medizinischen Bereich und fügt weitere zwei Rollen hinzu, nämlich *welcomer* und *family support*, die unweigerlich aufgrund von nicht vorhandenen Empfangs- und Entlassungsstrukturen in Einrichtungen des Gesundheitswesens entstanden sind. DolmetscherInnen sind nicht lediglich bei Gesprächen zwischen medizinischem Personal und PatientInnen dabei, sie begrüßen die PatientInnen unmittelbar vor der Dolmetschung, informieren sie außerhalb des Krankenhauses über institutionelle Gegebenheiten, erledigen Terminvereinbarungen und begleiten sie sogar zu Untersuchungen (vgl. Havelka 2018: 82f, Leanza 2007: 21). Es lässt sich also feststellen, dass das Dolmetschen im Gesundheitswesen sogar für berufliche DolmetscherInnen eine sehr herausfordernde und komplexe Aufgabe ist. Wie schaut die Situation aus, wenn – wie in den meisten Fällen – LaiendolmetscherInnen eingesetzt werden?

Kommen z.B. Reinigungskräfte mit Migrationshintergrund zum Einsatz, besteht die Gefahr, dass sie im Fall mangelnden medizinischen Hintergrundwissens improvisieren, was natürlich den Behandlungsprozess stark beeinflussen kann. Bei dolmetschenden Familienmitgliedern besteht nicht nur das Problem des fehlenden Fachwissens; auch die starke Eingebundenheit in die Familienhierarchie kann sich auf die Professionalität des Gesprächs auswirken. Genauso, wie man einem Arzt nicht widerspricht, verbietet es sich, ein „höher“ stehendes Familienmitglied zu unterbrechen oder zu korrigieren. Vor allem dann, wenn Kinder als DolmetscherInnen eingesetzt werden, können „Tabuthemen“ auftauchen, die zu Schamgefühl oder zu Fehlinterpretationen führen können. Auch mehrsprachige KrankenpflegerInnen kommen oft zum Einsatz. Sie verfügen zwar über das medizinische Fachwissen, müssen aber neben ihren eigentlichen beruflichen Aufgaben eine zusätzliche Arbeitsbelastung auf sich nehmen, wofür sie oft nicht einmal bezahlt werden. Der Mangel an professionellen Dolmetschkompetenzen, wie Unparteilichkeit, Macht-Einschätzung oder Neutralität, kann in bestimmten Fällen sogar zu Loyalitätskonflikten zwischen ArbeitgeberInnen und „Landsleuten“ führen (vgl. Hoefert & Härter 2010: 154).

5.2 Dolmetschen im österreichischen Gesundheitswesen

Es taucht also die berechtigte Frage auf: Warum werden in einem solch komplexen und wichtigen Bereich wie im Gesundheitswesen nicht professionelle DolmetscherInnen eingesetzt, die in der Lage wären, all diese Missverständnisse zu vermeiden?

Die Frage lässt sich aus verschiedenen Blickwinkeln beantworten. Aus Sicht des medizinischen Fachpersonals kann das Fehlen von gesetzlichen Regelungen als ein wesentlicher Faktor identifiziert werden. Dies bezieht sich sowohl auf die Regelungen über die Notwendigkeit des Einsatzes von DolmetscherInnen als auch auf die Finanzierung der Dolmetschleistung. Suchen wir nach Regelungen in Österreich, tauchen oft unverbindliche Worte wie „empfehlenswert“, „sollte man“ oder „ratsam“ auf (vgl. Straub 2016: 14f). Es wird nicht über DolmetscherInnen, sondern über „SprachmittlerInnen“ gesprochen, was bedeutet, dass „jede ausreichend sprachkundige Person“ (Straub 2016: 14) in Betracht gezogen werden kann. Die Tatsache, dass die Verständigungsschwierigkeiten zwischen ÄrztInnen und PatientInnen sehr häufig ein Problem darstellen, wird zwar erkannt, im Sozialversicherungsrecht gibt es aber keine Regelung darüber, in welchen Fällen DolmetscherInnen eingesetzt werden müssen und wer die Dolmetschkosten im Rahmen einer ärztlichen Behandlung zu tragen hat. Derartige Kosten stellen keine vertragsärztliche Leistung dar und können somit auch nicht an den Krankenversicherungsträger weiterverrechnet werden (vgl. Straub 2016: 15).

Aus Sicht der ausgebildeten DolmetscherInnen muss auch ein anderer Faktor erwähnt werden, der für den Mangel an qualitativer Sprachmittlung im Gesundheitsbereich verantwortlich sein kann: die Frage des Prestiges. Dolmetschen im Gesundheitswesen wird in der Literatur meist zum Bereich Community Interpreting gezählt. Community Interpreter (CI), oft auch als KommunaldolmetscherInnen bezeichnet, „ermöglichen Menschen, deren Mutter- bzw. Bildungssprache nicht die des Gastlandes ist, den Zugang zu sozialen und kommunalen Einrichtungen⁷ des Gastlandes.“ (Pöllabauer 2005: 53) Es lässt sich allgemein feststellen, dass das Prestige der CI nicht sehr hoch – und keinesfalls mit dem Ansehen der KonferenzdolmetscherInnen zu vergleichen ist. Auch wenn der Arbeitsbereich des Kommunaldolmetschens sehr abwechslungsreich ist und zahlreiche Herausforderungen bereithält, werden hier oft Ad-hoc-DolmetscherInnen teils ohne Bezahlung eingesetzt, deswegen müssen auch ausgebildete DolmetscherInnen mit einer niedrigeren Vergütung rechnen. Dieses Betätigungsfeld wird somit für qualifizierte TranslatorInnen oftmals ökonomisch uninteressant (vgl. Sauerwein 2007: 10f).

Es gab Initiativen und Pilotprojekte in Österreich, um den DolmetscherInnenmangel im Gesundheitsbereich zu beseitigen. Dank einer Kooperation zwischen der österreichischen Plattform Patientensicherheit und dem Bundesministerium für Gesundheit wurde im Mai 2011 die Arbeitsgruppe „Umgang mit nicht-deutsch-sprachigen Patienten“ gegründet. Noch im selben Jahr wurde das Pilotprojekt „Videodolmetschen im Gesundheitswesen“ im Rahmen der

⁷ Zu solchen Einrichtungen zählen auch Krankenhäuser.

wissenschaftlichen Tagung „Wie viel Deutsch braucht man, um gesund zu sein? Migration, Gesundheit und Übersetzung“ ins Leben gerufen. An der Entwicklung des Konzeptes nahmen mehrere Institutionen teil, unter anderem das Institut für Ethik und Recht in der Medizin und das Zentrum für Translationswissenschaft der Universität Wien. Das Projekt begann am 7. Oktober 2013 in zwölf Krankenhäusern und zehn Ärztezentren.⁸ Nach der erfolgreichen Testphase wurde im März 2014 die SAVD Videodolmetschen GmbH gegründet. Es muss hier jedoch erwähnt werden, dass im Rahmen dieses Videodolmetschdienst-Projektes damals nur die Sprachen Türkisch, Bosnisch/Serbisch/Kroatisch und die Gebärdensprache angeboten wurden. Auch die Zahl der DolmetscherInnen, die von einem Büro in Wien aus arbeiteten, war sehr beschränkt.⁹ Im Laufe der Jahre wuchs das Sprachangebot der SAVD Videodolmetschen GmbH: Mit mehr als 50 Sprachen zählt der Sprachdienstleister derzeit zu den Marktführern für Videodolmetschen im deutschsprachigen Raum. SAVD verspricht die direkte Zuschaltung von qualifizierten DolmetscherInnen in 14 Sprachen (Albanisch, Arabisch – Hocharabisch, Bulgarisch, Dari, Farsi/Persisch, Russisch, Kurdisch – Kurmanci, Serbisch, Polnisch, Slowakisch, Rumänisch, Tschechisch, Türkisch und Ungarisch) innerhalb von höchstens 4 Minuten; weitere 40 Sprachen können als Termin-Dolmetschungen gebucht werden. Die DolmetscherInnen stehen an Werktagen zwischen 08:00 und 18:00 Uhr in allen 54 Sprachen zur Verfügung; zwischen 18:00 und 08:00 Uhr sowie an Wochenenden und Feiertagen werden lediglich die vorher erwähnten 14 Sprachen – natürlich mit viel längeren Wartezeiten – angeboten (vgl. SAVD 2020, SAVD Videodolmetschen 2020). Dieses Sprachenangebot ist aber nicht ausreichend, allein wenn man bedenkt, dass in Wien ca. 2 Millionen Menschen aus 183 verschiedenen Ländern wohnen.¹⁰ Die Nachfrage mit ausgebildeten DolmetscherInnen in so vielen unterschiedlichen Sprachen vollständig abzudecken wäre allerdings eine fast unmögliche Aufgabe, vor allem weil diesbezüglich auch die finanziellen und rechtlichen Regelungen fehlen.

Dank der Entwicklung maschineller Übersetzungs- und Dolmetsch-Systeme gibt es heutzutage auch eine andere, kostengünstigere Lösung, die Sprachbarrieren im Krankenhaus bzw. in der Ordination zu überwinden. SST-Systeme werden immer öfter eingesetzt, um die Kommunikation auch ohne menschliche Dolmetschleistung zu ermöglichen. Dies bedeutet, dass heutzutage die Kommunikation allein mit Hilfe einer Applikation auf dem Handy zustande

⁸ Unter anderem zwei Spitäler des Krankenanstaltenverbundes (Sammelweis-Klinik, Rudolfstiftung), das St. Anna Kinderspital, das Lorenz-Böhler-Krankenhaus und das Meidlinger Unfallkrankenhaus (vgl. ORF 2013).

⁹ 8-14 DolmetscherInnen, 16 Stunden am Tag (von 6.00 Uhr bis 22.00 Uhr) (vgl. ORF 2013).

¹⁰ Ungarn ist die neuntgrößte Nationsgruppe in Wien (22.287 Menschen) nach Österreich, Serbien, Türkei, Deutschland, Polen, Rumänien, Bosnien und Kroatien. (vgl. Wetz 2016; Stadt Wien 2020)

kommen kann. Die berechtigte Frage lautet aber: Inwiefern sind diese Apps heute schon dazu geeignet, HumandolmetscherInnen zu ersetzen und eine erfolgreiche Kommunikation zwischen Arzt/Ärztin und PatientIn zu ermöglichen?

6. Forschungsstand

In diesem Kapitel werden einige Untersuchungen, die bereits mit der Dolmetsch-Applikation *SayHi Translate* durchgeführt wurden, vorgestellt.

Brizar (2014) stellte in ihrer Arbeit einen funktionalen Vergleich zwischen drei SST-Applikationen, *Jibbig Translator*, *VoiceTra4U* und *SayHi Translate*, an. Um die Stärken und Schwächen der einzelnen Applikationen bezüglich des Funktionsumfanges zu präsentieren, wurden SWOT-Analysen (engl. Akronym für Strengths (Stärken), Weaknesses (Schwächen), Opportunities (Chancen) und Threats (Risiken)) erstellt. Bezüglich der Stärken von *SayHi Translate* wurden insbesondere die große Anzahl der Sprachen und Dialekte, die Robustheit bzw. Verlässlichkeit der App (keine Abstürze oder technische Störungen) und der transparente Verlauf der Sprachverarbeitung hervorgehoben (vgl. Brizar 2014: 101). Zum unmittelbaren Vergleich der Übersetzungsleistungen der verschiedenen Applikationen wurden Testdialoge aus jeweils zwei Kontexten (touristisch und geschäftlich) und in verschiedenen Sprachkombinationen (DE-EN, DE-FR) durchgeführt. Bei diesen Gesprächen wurden bewusst Charakteristika der Spontansprache (z.B. Denkpausen, Selbstkorrekturen, syntaktisch falsche Textsegmente usw.) eingefügt (vgl. 2014: 103), was darauf schließen lässt, dass die Texte der Testdialoge zwar spontan schienen, jedoch mehr oder weniger vorgeschrieben wurden. Bei der Analyse der Testdialoge betonte Brizar die positive Leistung der Spracherkennungskomponente, die trotz einiger Fehler deutlich bessere Transkriptionen lieferte als die der beiden anderen Apps. Außerdem wird auch die gute Leistung der maschinellen Übersetzungskomponente hervorgehoben, die laut der Autorin „durchaus zufriedenstellende Ergebnisse“ (2014: 114) liefert. Darüber hinaus wird positiv angemerkt, dass „SayHi als einzige App Rücksicht auf Groß- und Kleinschreibung nimmt.“ (2014: 114) Brizar gelangt in ihrer Arbeit zu dem Schluss, dass die untersuchten SST-Applikationen für alltägliche Szenarien (z.B. Reisen) durchaus anwendbar seien – da Fehler in einem solchen Kontext leichter zu verzeihen sind –, ihre Anwendung in fachspezifischen Bereichen jedoch nicht ratsam wäre (vgl. 2014: 122).

Eine weitere Arbeit, die im Bereich des maschinellen Dolmetschens erstellt wurde und die Applikation *SayHi Translate* untersucht, ist jene von Mitrovic (2017). Mitrovic vergleicht in ihrer Arbeit die Leistung von zwei mobilen Apps, nämlich *Google Translate* und *SayHi Translate*, die im Bereich des Gesprächdolmetschens im Sprachenpaar Bosnisch/Kroatisch/Serbisch und Deutsch untersucht wurden. Im Rahmen der Versuche wurden die Apps im Hinblick auf Alltagssprache, Dialekt und medizinische Fachsprache betrachtet. Auch sie stellte fest, dass die Spracherkennung der App durchaus gut funktioniert, die gesprochenen Sätze

können jedoch ihrer Meinung nach „nicht richtig analysiert und übersetzt bzw. gedolmetscht werden“. (Mitrovic 2017: 119) Ihre Untersuchungen zeigen, dass sowohl allgemeinsprachliche als auch fachsprachliche Ausdrücke ein Problem für das System darstellen können und dass die Regeln der Syntax und Grammatik oft nicht beachtet werden. Bezüglich der Leistung der Spracherkennungskomponente betont Mitrovic, dass die Sätze unter günstigen Bedingungen gesprochen wurden, und stellt sich die Frage, „wie die Sprachverarbeitung beispielsweise mit Hintergrundgeräuschen aussehen würde.“ (2017: 119) Mitrovic weist in ihrer Arbeit auch auf die unangenehme, synthetische Stimme der Sprachausgabe hin. Sie kam anhand ihrer Untersuchungen zu dem Schluss, dass zwar die Übertragung einfacher Sätze mit den Smartphone-Apps grundsätzlich möglich ist, jedoch „weder *Google Translate* noch *SayHi* für die Verarbeitung gesprochener Sequenzen aus dem medizinischen Bereich für die Sprachrichtungen Bosnisch, Kroatisch, Serbisch – Deutsch und umgekehrt geeignet sind“. (2017: 119) Da längere, fachsprachliche Textpassagen nicht einwandfrei von den Applikationen übersetzt werden können, kann laut Mitrovic die Anwendung solcher Apps im medizinischen Bereich sogar als gefährlich eingestuft werden (vgl. 2017: 120). Sie erwähnt in ihrer Arbeit einen weiteren interessanten Punkt, nämlich dass SST-Apps einige Aspekte „wie beispielsweise kulturelle Unterschiede oder die nonverbale Kommunikation“ (2017: 120), die aus Sicht des Dolmetschens essenziell wären, nicht berücksichtigen können (vgl. 2017: 120).

Auch die Arbeit von Leitner (2020) darf nicht unerwähnt bleiben, in der eine vergleichende Analyse zwischen Mensch und Maschine auf Basis eines Experimentes mit der MD-Applikation *iTranslate Converse* im Sprachenpaar Deutsch-Französisch durchgeführt wurde. Anders als in den bereits erwähnten Arbeiten handelte es sich in diesem Fall um einen längeren, sehr realitätsnahen Dialog, der im Rahmen einer Vernehmung zwischen einem deutschsprachigen Polizeibeamten und einer französischsprachigen Person, die Opfer eines Diebstahls geworden war, stattfand. Beim ersten Durchgang wurde die sprachübergreifende Kommunikation mit Hilfe der App *iTranslate Converse* und im zweiten Schritt durch eine Humandolmetscherin durchgeführt (vgl. Leitner 2020: 55). Obwohl Leitner sich in ihrer Arbeit nicht mit der Applikation *SayHi Translate* beschäftigt, fand ich ihre Überlegungen in Hinblick auf das Experiment innovativ und zielführend, weil „es sich nicht um einen vorgeschriebenen und von den GesprächspartnerInnen abgelesenen Dialog handelte, sondern die Kommunikation zwischen den beiden Personen spontan erfolgte.“ (2020: 56) Dies wird auch in der vorliegenden Arbeit der Fall sein. Das Gespräch wurde auf Video aufgenommen, um die Mimik und Gestik aller TeilnehmerInnen erfassen zu können (vgl. 2020: 56). Dies ist auch meiner Meinung nach für die Analyse von großer Bedeutung, da die Körpersprache der Kommunikationsparteien sehr viel

darüber verraten kann, ob die Übersetzung bzw. Dolmetschung tatsächlich verstanden wurde. Leitner schließt aus dem durchgeführten Experiment und der anschließend erfolgten vergleichenden Analyse, dass die Dolmetsch-App „zu viele Fehlerquellen aufweist, als dass sie im kommunalen Bereich Anwendung finden könnte.“ (2020: 112) Auch sie betont, dass solche Applikationen für den alltäglichen Gebrauch eine Hilfe darstellen können, bezweifelt jedoch, dass *iTranslate Converse* „eines Tages den Beruf von KommunaldolmetscherInnen gefährden könnte.“ (2020: 112)

Eine weitere Studie, die sich mit der Verwendbarkeit des maschinellen Übersetzungssystems *Google Translate* im medizinischen Bereich beschäftigt, wurde 2014 im Vereinigten Königreich durchgeführt. Im Rahmen dieser Untersuchung, deren Beschreibung und Ergebnisse unter dem Titel „Use of Google Translate in medical communication: evaluation of accuracy“ im *British Medical Journal* publiziert wurden, wurden jeweils 10 medizinische Sätze der AS Englisch mittels *Google Translate* in 26 Sprachen (8 westeuropäische, 5 osteuropäische, 11 asiatische und 2 afrikanische) übersetzt, was insgesamt 260 übersetzte Sätze ergab. Die Übersetzungen wurden dann von MuttersprachlerInnen der gegebenen Sprachen zurück ins Englische transferiert, um zu überprüfen, inwiefern diese vom Originalsatz abweichen. Bei der Evaluierung der Übersetzungen waren Sätze, die keinen Sinn ergaben oder sachlich inkorrekt waren, als falsch eingestuft; kleinere grammatikalische Fehler waren erlaubt. Von den gesamten Übersetzungen waren 150 (57,7%) korrekt und 110 (42,3%) falsch. Afrikanische Sprachen erzielten die niedrigsten Werte (45% korrekt), gefolgt von asiatischen Sprachen (46%). Osteuropäische Sprachen erreichten 62%, während der Anteil der als korrekt eingestuften Übersetzungen bei den westeuropäischen Sprachen mit 74% am größten war (vgl. Patil & Davies 2014).

Es gab einige schwerwiegende Fehler. Zum Beispiel wurde „Your child is fitting“ auf Suaheli als „Your child is dead“ übersetzt. Der Satz „Your husband has the opportunity to donate his organs“ wurde als „Your husband can donate his tools“ ins Polnische übertragen (vgl. Patil & Davies 2014). Die Liste der medizinischen Sätze, die im Rahmen dieser Untersuchung übersetzt wurden, sind in Tabelle 3 angeführt.

Tabelle 3: Liste der medizinischen Sätze, die mit Google Translate übersetzt wurden (Patil & Davies 2014)

Phrase translated	Sample or most common error	Percentage correct
Your wife is stable	Your wife cannot fall over	53.8
Your husband had a cardiac arrest	Your husband's heart was imprisoned	53.8
Your husband had a heart attack	Your husband's heart was attacked	73.1
Your wife needs to be ventilated	Your wife needs to be aired	26.9
Your child's condition is life threatening	Your child's state is not life stopping	69.2
Your child has been fitting	Your child has been constructing	7.7
Your child will be born premature	Your child is sleeping early	76.9
Your husband has the opportunity to donate his organs	Your husband is now ready to donate	88.5
We will need your consent for operation	We need your consent for operating (such as machinery)	61.5
Did he have high fever at home?	Your home temperature was high	65.4

Google Translate konnte insgesamt nur 57,7% der Sätze korrekt übersetzen und ist laut BMJ im Rahmen medizinischer Gespräche auf keinen Fall vertrauenswürdig. Im Bericht wird betont, dass, obwohl die Verwendung solcher Applikationen oft als die einfachste und kostengünstigste Lösung für die Überwindung der Sprachbarrieren erscheint, sie nicht bei ernsteren Beschwerden oder lebensrettender ärztlicher Kommunikation verwendet werden dürfen (vgl. Patil & Davies 2014).

Eine ähnliche Studie wurde auch in Australien im April 2019 durchgeführt, hier wurden aber insgesamt 10 mobile Applikationen im Krankenhausalltag getestet (*SayHi Translate*, *Google Translate*, *Microsoft Translator*, *MediBabble Translator*, *Universal Doctor Speaker*, *Canopy Speak*, *TripLingo*, *Naver Papago Translate*, *CALD Assist*, *Talk To Me*). Die Studie wurde unter dem Titel „Language Translation Apps in Health Care Settings: Expert Opinion“ auf der Webseite JMIR Publications veröffentlicht (vgl. Panayiotou et al. 2019). Ziel dieser Untersuchung war es, iPad-kompatible Übersetzungs-Apps zu evaluieren, um beurteilen zu können, ob sie für alltägliche Gespräche im Gesundheitswesen geeignet sind. Es wurden nur Applikationen ausgewählt, die kostenlos verfügbar waren und mindestens eine der zehn in Australien meist gesprochenen Sprachen übersetzen konnten. Die Übersetzungen erfolgten immer zwischen dem Englischen und einer der folgenden Sprachen: Mandarin, Arabisch, Kantonisch, Vietnamesisch, Italienisch, Griechisch, Hindi, Spanisch, Panjabi und Tagalog. Im ersten Schritt wurde die BenutzerInnenfreundlichkeit der Apps untersucht (Offline-Verwendung,

Spracheingabe und -ausgabe, Anzahl der Sprachen). Im zweiten Schritt der Untersuchung wurde die Tauglichkeit der Apps für alltägliche Gespräche im Gesundheitswesen von zwei unabhängigen ExpertInnen interkultureller Kommunikation analysiert. *SayHi Translate* erwies sich im Rahmen dieser Studie als wenig geeignet für den Einsatz im Bereich der medizinischen Kommunikation (vgl. Panayiotou et al. 2019).

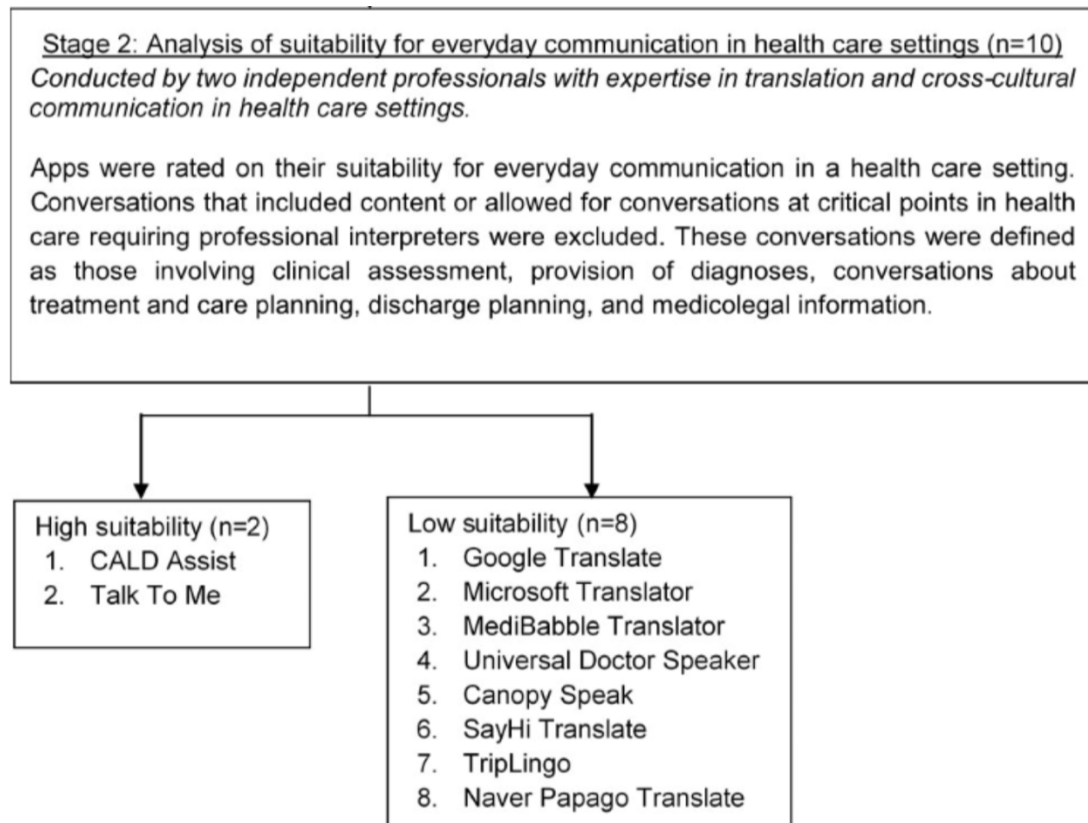


Abbildung 19: Gesamtbewertung der Apps bezüglich der Tauglichkeit für alltägliche Gespräche im Gesundheitswesen (Panayiotou et al. 2019)

7. Experiment

Wie bereits erwähnt, bezieht sich die zentrale Fragestellung dieser Masterarbeit auf die Evaluierung einer Arzt-Patient-Kommunikation im Sprachenpaar Deutsch-Ungarisch, die lediglich mit Hilfe der App *SayHi Translate* zustande kam. In Kapitel 6 wurden die bereits durchgeführten Untersuchungen mit Dolmetsch-Applikationen erläutert. Die meisten Experimente konzentrieren sich einerseits auf die BenutzerInnenfreundlichkeit, andererseits auf die Evaluierung der Übersetzungsqualität. Die Analyse wurde dabei oft auf Basis lexikalischer bzw. grammatikalischer Korrektheit, Verständlichkeit, Stil, Kohärenz, Kohäsion und Informationstransfer der Übersetzung erstellt. Diese Faktoren, die bei der Untersuchung einer Übersetzungsleistung sehr relevant sind, bilden nur einen Teil meiner Evaluierungskriterien. Ziel meines Experimentes war es, eine möglichst realitätsnahe Situation zu schaffen, in der auch andere Faktoren untersucht werden können, die im wirklichen Leben eine wesentliche Rolle spielen: Neben der Leistung von *SayHi Translate* wurden auch Aspekte behandelt, die sich darauf beziehen, inwiefern die Verwendung eines Tablets als „Kommunikationsbrücke“ während einer ärztlichen Untersuchung die Natürlichkeit der menschlichen Interaktion beeinflusst oder verändert. Nach welchen Aspekten die Evaluierung durchgeführt wurde, wird in Kapitel 7.4 detailliert vorgestellt.

7.1 Forschungsfragen

Die Forschungsfragen, die ich anhand des Experimentes beantworten möchte, sind folgende:

- Inwieweit ist *SayHi Translate* dafür geeignet, die Kommunikation im Rahmen des Arzt-Patient-Gesprächs im Sprachenpaar Deutsch-Ungarisch zu ermöglichen?
- Waren die Teilnehmer des Experimentes mit der Leistung von *SayHi Translate* zufrieden? Würden sie auf Basis ihrer Erfahrungen auf diese Applikation auch in der Zukunft in einer ähnlichen Situation gerne zurückgreifen?

7.2 Teilnehmer des Experiments

Das Experiment fand am 7. Dezember 2019 in Wien in einer Ordination statt. Die Rolle des Arztes wurde von einem in Wien geborenen österreichischen Staatsbürger eingenommen, der als Facharzt für Chirurgie und Gefäßchirurgie tätig ist. Da sich in seinem Patientenkreis

Menschen aus verschiedenen Ländern befinden (z.B. Russland), arbeitet er oft mit DolmetscherInnen zusammen.

Die Person, die im Rahmen des Experimentes die Rolle des Patienten einnahm, ist ein in Ungarn lebender, ungarischer Staatsbürger, der jahrelang als Pädagoge tätig war und derzeit im ordnungsamtlichen Bereich arbeitet.

Wie bereits erwähnt, war die Zielsetzung des Experiments eine möglichst realitätsnahe Situation zu schaffen. Aus diesem Grund war es mir wichtig, die TeilnehmerInnen so auszuwählen, dass sie keine gemeinsame Sprache besitzen. Der Arzt spricht fließend Deutsch und kann sich auch auf Englisch sehr gut verständigen, versteht aber kein Wort Ungarisch. Der Patient ist der deutschen Sprache gar nicht mächtig; außerdem sind seine Englisch-Grundkenntnisse nicht genügend, um ein medizinisches Gespräch zu führen. Er sprach während der Untersuchung ausschließlich in seiner Muttersprache, also Ungarisch. Auf diese Weise war die verbale Kommunikation ausschließlich auf die Leistung der App angewiesen. Hätte der ungarische Patient die Fragen des Arztes auf Deutsch mehr oder weniger verstanden, hätte er auch dann antworten können, wenn die Informationen von der App nicht vollständig oder richtig ins Ungarische übersetzt worden wären.

7.3 Ablauf der Untersuchung

Die Untersuchung war vollkommen der Wirklichkeit entsprechend: Der Patient wollte sich aufgrund seiner tatsächlichen Beschwerden untersuchen lassen, um herauszufinden, ob er – wie viele in seiner Familie – eine Neigung zu Krampfadern hat. Da sich seine Beine oft schwer und müde anfühlten, wollte er sich einer ärztlichen Untersuchung unterziehen, um der Krankheit vorzubeugen oder – falls nötig – sich behandeln zu lassen. Er hatte die App *SayHi Translate* schon einen Monat vor dem Experiment sowohl auf sein Smartphone als auch auf sein Tablet heruntergeladen, um sich mit der Verwendung der App vertraut zu machen.

Die Untersuchung war spontan, der Arzt und der Patient begegneten sich zum ersten Mal. Mit dem Patienten war nur besprochen worden, dass er am Anfang (mittels der App) kurz erklärt, wie *SayHi Translate* funktioniert, vor allem weil er die ungarische Version der App heruntergeladen hatte und so dem Arzt vielleicht nicht alle Funktionen bekannt und verständlich waren (z.B. wie man den Text manuell korrigieren kann).

Nach einer kurzen Begrüßung – bei der das Tablet noch nicht verwendet wurde – nahm der Patient am Schreibtisch Platz und informierte den Arzt kurz darüber, dass er wegen seiner

mangelnden Sprachkenntnisse eine Dolmetsch-App, nämlich *SayHi Translate*, auf seinem Tablet mitgebracht habe, um die zwischensprachliche Kommunikation zu ermöglichen (dies erfolgte schon mit Hilfe der App). Der Patient erklärte dem Arzt die Nutzung der App in ein paar Sätzen und führte anschließend seine Beschwerden und Sorgen bezüglich seiner Krampfadern aus. Der Arzt stellte ihm ein paar Fragen (z.B. ob es sich um das linke oder das rechte Bein handelt; wann und wo die Beschwerden auftreten usw.) und bat ihn, seine Beine freizumachen und sich auf das kleine Podest zu stellen, um eine Ultraschalluntersuchung durchführen zu können. Die Untersuchung dauerte ungefähr 4 Minuten. Nach der Untersuchung bat der Arzt den Patienten sich wieder anzuziehen und sich an den Schreibtisch zu setzen, um ihm das Problem an den Beinvenen anhand einer Grafik erklären zu können. Die Ergebnisse der Untersuchung waren beruhigend. Der Patient hat das Problem, an dem sein Großvater und Vater litten, glücklicherweise nicht geerbt, denn es wurde anhand der Ultraschalluntersuchung diagnostiziert, dass alle Venenklappen sehr gut funktionieren. Er erhielt vom Arzt einige Tipps für eine aktive und gesunde Lebensführung, um Krampfader-Problemen vorzubeugen, sowie ein Rezept für ein Pflanzen-Extrakt-Medikament, das in der Prävention helfen kann. Der Patient bedankte sich für die Untersuchung und die hilfreichen Ratschläge und sie verabschiedeten sich.

Die Untersuchung dauerte 28 Minuten und wurde auf Video aufgenommen¹¹. Auf dem verwendeten Tablet wurde auch eine Bildschirm-Aufnahme gemacht, um die Textdaten bei der Analyse genauer beobachten zu können. Nach dem Experiment wurde sowohl mit dem Arzt als auch mit dem Patienten ein Interview durchgeführt.

7.4 Evaluierungsmethode

7.4.1 Problematik der Evaluierung von MD-Systemen

Möchte man eine Evaluierungsmethode für die Beurteilung und Bewertung eines maschinellen Dolmetschsystems festlegen, ist man schnell mit verschiedenen Problemen konfrontiert. Es gibt nämlich keine allgemein anwendbare, universell akzeptierte Evaluationsmethode oder Maßstäbe, die bei allen Systemen verwendet werden könnten. Es fehlen außerdem allgemein gültige Qualitätsstandards und Qualitätsnormen für die Evaluierung von Übersetzungs- bzw. Dolmetsch-Apps. Der Großteil der Methoden, mit denen solche Analysen durchgeführt werden, ist anwendungsabhängig und kann nur für das ausgewählte System verwendet werden. Eine

¹¹ Ich möchte an dieser Stelle anmerken, dass nach 24 Minuten und 36 Sekunden der Speicher meines iPhones voll wurde und deswegen die letzten 3 Minuten mit einem anderen Handy aufgenommen wurden. Dies erklärt den Bruch im Videomaterial.

weitere Herausforderung bei der Untersuchung der Übersetzungsleistung ist in der Natur der Übersetzung zu finden: Man kann faktisch nicht von einer „korrekten“ Übersetzung sprechen, es gibt in den meisten Fällen mehrere entsprechende Übersetzungsmöglichkeiten. Es darf außerdem nicht vergessen werden, dass BenutzerInnen, EntwicklerInnen und ForscherInnen unterschiedliche Sichtweisen und Vorstellungen hinsichtlich der Evaluierung haben, was unweigerlich zu Kontroversen führt. Cap (2003) befasst sich unter anderem mit der Problematik der Evaluierung von MÜ-Systemen und stellt Folgendes fest:

Da eine Vielzahl an Evaluierungen die Qualität der Übersetzungen beziehungsweise der Übersetzungssysteme nicht komplett erfasst und meist (aus Zeit- und Kostengründen) nur einen Ausschnitt des Programms untersucht oder insgesamt eine eher oberflächliche Analyse liefert und sogar Fehler beziehungsweise Ungenauigkeiten während des Evaluierungsvorgangs einräumt, steht der Neuling, aber auch der erfahrene Evaluator vor der Wahl, entweder die bestehenden Evaluierungspraktiken zu verbessern oder neue und bessere Evaluierungsmethoden zu entwickeln. (Cap 2003: 46)

Bei jeglicher Evaluierung sind zwei Faktoren ausschlaggebend: Zum einen soll das Objekt unvoreingenommen, objektiv und fair untersucht werden, das heißt, dass während der Evaluation eventuelle Vorurteile keine Rolle spielen dürfen. Zum anderen müssen die Ergebnisse der Evaluierung relevant und informativ sein. Um dies realisieren zu können, muss man das getestete Programm gut kennen und die zu messenden Größen präzise erfassen. In welcher Form die Evaluierung strukturiert sein soll, sollte von Beginn an klar festgelegt werden (vgl. Cap 2003: 46). Da die Brauchbarkeit der Übersetzungen und der gesamten App von den Anforderungen und Bedürfnissen der BenutzerInnen abhängt, wurde die Evaluierungsmethode so zusammengesetzt, dass neben den linguistischen Kriterien auch benutzerorientierte Aspekte berücksichtigt werden. Die Analyse erfolgt in einem spezifizierten Kontext, in einer speziellen Arbeitssituation. Um aussagekräftige Ergebnisse zu bekommen, habe ich mich entschlossen, aufgabenbezogen vorzugehen, das heißt, dass „die Bewertung eines Systems in Abhängigkeit vom angestrebten Einsatzgebiet“ (Cap 2003: 50) erfolgt. Im folgenden Unterkapitel wird die angewendete aufgabenbezogene Evaluierungsmethode genauer vorgestellt.

7.4.2 Aufgabenbezogene Evaluierung

Für lange Zeit waren für die Bewertung maschineller Übersetzungssysteme lediglich die Art und Anzahl von Fehlern und Kriterien wie Informationsgehalt oder Verständlichkeit der übersetzten Texte ausschlaggebend. Heutzutage gewinnen aber aufgaben- und benutzerspezifische Faktoren immer mehr an Bedeutung (vgl. Cap 2003: 50), vor allem bei maschinellen

Dolmetschsystemen, die oft in konkreten Gesprächssituationen als „Hilfe“ verwendet werden, um etwas Bestimmtes herauszufinden (z.B. wo es eine Apotheke in der Nähe gibt, ob der Kellner uns ein Gericht empfehlen kann, usw.). An dieser Stelle beziehe ich mich auf die Skopos-theorie (Skopos (gr.) = „Zweck“, „Ziel“, „Absicht“) von Hans J. Vermeer und Katharina Reiß, bei der das Handeln im Mittelpunkt steht, das immer durch eine Situation und eine Intention bestimmt wird. Der Handelnde ist sich der Situation bewusst und verfolgt mit seinem Tun immer einen konkreten Zweck. Das bedeutet, dass jede Handlung von einer bestimmten Motivation getrieben wird (vgl. Reiß/Vermeer 1991). Da Translation als eine Unterkategorie translatorischen Handelns aufgefasst werden kann (vgl. Vermeer 1990: 94), sind der Zweck und die Funktion der Übersetzung bei der Evaluierung maschineller Übersetzungs- und Dolmetschsysteme nicht wegzudenken. Dementsprechend ist auf dem Gebiet der Sprachverarbeitung eine Entwicklung von genereller Evaluierung hin zu aufgaben- und zweckorientierter Evaluierung zu beobachten (vgl. Cap 2003: 50).

Wie bereits in Kapitel 5 dargestellt, verfolgen während der Untersuchung sowohl der Patient als auch der Arzt ein bestimmtes Kommunikationsziel: Der Patient erwartet in dieser Situation ärztliches Verständnis und Gehör für seine Beschwerden und möchte herausfinden, ob seine Sorgen bezüglich der Krampfadern berechtigt sind; der Arzt möchte inzwischen relevante Informationen für die Erstellung einer Diagnose bekommen, die Ergebnisse dem Patienten erklären und ihn – falls nötig – über die weiteren medizinischen Schritte informieren. Für eine gelungene Kommunikation ist die Zufriedenheit beider Kommunikationsparteien unerlässlich. Es kann vorkommen, dass die App einige Aussagen grammatikalisch oder stilistisch falsch übersetzt – werden diese jedoch von den Experimentteilnehmern gut verstanden, ist das Kommunikationsziel grundsätzlich erreicht.

Natürlich bedeutet das nicht, dass linguistische Kriterien bei der Evaluierung nicht wichtig wären: Lexikalische Korrektheit, Kohäsion, Kohärenz, Grammatik usw. sind sehr relevante Aspekte, auf die der Informationstransfer aufbaut. Aus diesem Grund bildet auch die linguistische Analyse ein wesentlicher Teil der Evaluierung.

Die von mir verwendete erweiterte aufgabenorientierte Evaluierungsmethode setzt sich aus vier Punkten zusammen. Zuerst wird die **BenutzerInnenfreundlichkeit** der App analysiert: In diesem Teil der Evaluierung wird die BenutzerInnenschnittstelle (User Interface), also Aufbau der App, Design usw., untersucht. Grundsätzlich befasst sich dieser Punkt mit der Frage: Wie einfach ist es, die App zu verwenden? Im zweiten Punkt werden die **Spracherkennungsfehler** beleuchtet. Wie schon in Kapitel 3.1.1 dargelegt, können verschiedene Störfaktoren, wie Dialekte, Akzente, Sprechgewohnheiten, Raumakustik usw. die automatische

Spracherkennung stark beeinflussen und zu Fehlern bei der Transkription führen. In diesem Teil werden die während des Experimentes aufgetauchten Worterkennungsfehler aufgelistet und ihre negative Wirkung auf den Informationstransfer analysiert. Der dritte Punkt der Evaluierung setzt sich mit der **linguistischen Analyse** auseinander. Hier werden die von der App übersetzten Sätze und Sequenzen nach Aspekten wie lexikalischer bzw. grammatikalischer Korrektheit, Kohärenz, Kohäsion, Stil, Wortstellungsfehler usw. untersucht. Schließlich wird das Experiment aus Sicht der **Usability** analysiert. In meiner Arbeit wird unter „Usability“ die „Natürlichkeit der Interaktion“ verstanden, also inwiefern die Tatsache, dass die Kommunikation mit Hilfe eines Handys/Tablets ermöglicht wird, die Interaktion zwischen zwei Menschen beeinflusst oder verändert: Wer hält das Tablet in der Hand? Schauen die Teilnehmer oft auf den Bildschirm? Bleibt der Blickkontakt trotz des Tablets erhalten? Wie lange braucht die App, um die Sequenzen zu „dolmetschen“? Muss man während der Untersuchung Sätze wiederholen oder manuell korrigieren, um Missverständnisse zu vermeiden? Kommt es wegen der zu geringen Geschwindigkeit von *SayHi Translate* zu längeren Pausen oder muss das Experiment wegen der schlechten Leistung der App unterbrochen oder sogar völlig abgebrochen werden?

Mit der Einführung dieses letzten Aspekts möchte ich die bereits bestehenden Evaluierungsmethoden, die sich meistens lediglich auf die Analyse der BenutzerInnenfreundlichkeit und der Übersetzungsleistung einer MD-App konzentrieren, um einen zusätzlichen Blickwinkel erweitern. Die Idee der Einführung dieses Punktes basiert auf dem Gedanken, dass während eines Gespräches – unabhängig davon, ob dieses innerhalb einer Sprache oder zwischen verschiedenen Sprachen geführt wird – viele Informationen non-verbal übermittelt werden. Aus diesem Grund fand ich es sinnvoll, auch das Verhalten der Teilnehmer während des Experiments zu beobachten.

7.5 Transkription des Experimentes

Der Ablauf der ärztlichen Untersuchung mit *SayHi Translate* wurde in Form einer Tabelle transkribiert:

Tabelle 4: Ausschnitt aus der Transkription des Experimentes (Zeilen 1 – 5)

	Gesprochene Eingabe	ASR + Transkription	MÜ	Bildbeschreibung, nonverbale Kommunikation
	A	B	C	D
1	Hallo!	x	x	<i>gibt die Hand</i>
2	Bitte, setzen Sie sich!	x	x	<i>schüttelt die Hand und zeigt auf den Stuhl</i>
3	Jó napot kívánok doktor úr, köszönöm szépen, hogy időt szánt rám <i>pont</i> . Bécsben dolgozom, de sajnos nem beszélek németül, ezért hoztam el ezt az applikációt, ami remélem segít nekünk a kommunikációban <i>pont</i> .	Napot kívánok doktor úr köszönöm szépen hogy időt szánt rám. Bécsben dolgozom, de sajnos nem beszélek németül ezért hoztam el ezt az applikációt amit remélem segít nekünk majd a kommunikációban.	Ich wünsche Ihnen einen Tag, Doktor. Ich arbeite in Wien, aber leider spreche ich kein Deutsch, deshalb habe ich diese App mitgebracht, die uns hoffentlich helfen wird, zu kommunizieren.	<i>setzt sich an den Schreibtisch und beginnt zu sprechen; der Arzt möchte gleich antworten und greift nach dem Tablet, der Patient erhebt jedoch den Zeigefinger, um zu signalisieren, dass er noch etwas sagen möchte</i>
4	Gyorsan elmagyarázom, hogy hogyan működik a program. Itt vannak a nyelvek, amikor beszélünk, ezeket a gombokat kell megnyomni és utána kell beszélni <i>pont</i> .	Gyorsan elmagyarázom hogy hogyan működik a program, itt vannak a nyelvek, amikor beszélünk ezeket a gombokat kell megnyomni és utána kell beszélni.	Ich erkläre schnell, wie das Programm funktioniert, hier sind die Sprachen, wenn wir über diese Tasten sprechen, die gedrückt werden und dann sprechen.	<i>zeigt während der Spracheingabe auf die Tasten; nachdem das Gerät übersetzt hat dreht der Arzt das Tablet zu sich</i>
5	Gut, ich glaube, das habe ich verstanden <i>Punkt</i> . Darf ich Sie fragen, was ich für Sie tun kann?	Ich glaube das habe ich verstanden. Darf ich Sie fragen was ich für Sie tun kann	Azt hiszem, megértettem. Megkérdezhetem, hogy mit tehetek érted	

Wie im Ausschnitt (Tabelle 4) veranschaulicht wird, wurde die Transkription in 4 Spalten geteilt.

In der ersten Spalte (*A*) befinden sich die gesprochenen Eingaben und auch jene sprachlichen Äußerungen, die während der ärztlichen Untersuchung von den Teilnehmern artikuliert, aber nicht von der Applikation übersetzt wurden. Bei den Aussagen, wo *SayHi Translate* nicht verwendet wurde, befindet sich in den Spalten *B* und *C* ein „x“. Da in der Spontansprache kaum Pausen oder Satzzeichen zu erkennen sind, müssen diese lautsprachlich hinzugefügt werden. Natürlich ist es mühsam und ungewöhnlich *Punkt, Komma, Rufzeichen* oder *Fragezeichen* zu sagen. Die Teilnehmer des Experimentes haben versucht, die Satzzeichen bei der gesprochenen Eingabe nicht zu vergessen¹². Diese werden in der ersten Spalte (*A*) als *Punkt, Komma, Rufzeichen* und *Fragezeichen* / *pont, vessző, felkiáltójel*, und *kérdőjel* dargestellt.

In der zweiten Spalte (*B*) ist die von der Spracherkennungskomponente erstellte Transkription zu lesen. Hierbei ist zu betonen, dass die Transkription der AS genau so dargestellt wird, wie sie in der App zu sehen war, auch wenn sie Spracherkennungsfehler beinhaltete.

In der dritten Spalte (*C*) ist die maschinelle Übersetzung des erkannten Textes dargestellt.

Da ich auch die „Natürlichkeit der Interaktion“ untersucht habe, war es sinnvoll, eine vierte Spalte zu erstellen (*D, Bildbeschreibung, nonverbale Kommunikation*), in der die aus Sicht der Interaktion und der Kommunikation wichtigsten nonverbalen Kommunikationselemente der Teilnehmer beschrieben werden. An dieser Stelle möchte ich anmerken, dass diese Spalte nur bei jenen Zeilen ausgefüllt wurde, wo die nonverbale Kommunikation aus Sicht der Evaluierung erwähnenswert war.

Die Untersuchung kann in 3 Abschnitte geteilt werden: die Besprechung der Beschwerden (Zeilen 1 – 16), die Untersuchung (Zeilen 17 – 34) und das Aufklärungsgespräch bzw. Besprechung der Ergebnisse (Zeilen 35 – 71). Die Farben von *SayHi Translate* wurden auch in die Tabelle übernommen: Die Aussagen der Patienten sind blau, während das vom Arzt Gesagte grün markiert wurde.

Ausdrücke, die falsch oder gar nicht vom System erkannt wurden, sind mit einem Unterstrich in den Spalten *B* und *C* markiert. Wörter, die manuell korrigiert wurden, sind doppelt unterstrichen.

¹² Dass man die Arbeit der App wesentlich erleichtern kann, indem man die Satzzeichen einfach selbst ausspricht, habe ich dem Arzt im Vorfeld des Experiments kurz erklärt. Der Patient wusste dies schon aus eigener Erfahrung.

Bei der Fehleranalyse werden in den Beispielzeilen jene Ausdrücke, die gerade diskutiert werden, mit Gelb markiert. Ihre Übersetzungen werden mit Blau hervorgehoben. Bei einigen Spracherkennungsfehlern werden die falsch flektierten Verben in der Übersetzung mit Grün gekennzeichnet.

Tabelle 5: Ausschnitt aus der Transkription des Experimentes, Beispiel für einen Spracherkennungsfehler (Zeilen 46 – 49)

	Gesprochene Eingabe	ASR + Transkription	MÜ	Bildbeschreibung, nonverbale Kommunikation
	A	B	C	D
46	Damit haben Sie das Problem <i>Komma</i> , dass ihr Vater und ihr Großvater offenbar gehabt haben, nicht <u>geerbt</u> .	Damit haben Sie das Problem, dass ihr Vater und ihr Großvater offenbar gehabt haben nicht <u>geärgert</u>	Ez azt jelenti, hogy a probléma, hogy apja és nagyapja láthatóan nem <u>bosszantott</u>	
47	Nicht <u>geerbt</u> .	Nicht <u>geärgert</u>	Nem <u>bosszús</u>	<i>erkennt, dass das Ende des Satzes nicht richtig transkribiert wurde und wiederholt die letzten zwei Wörter</i>
48				<i>zeigt "es spinnt"</i>
49		Damit haben Sie das Problem, dass ihr Vater und ihr Großvater offenbar gehabt haben nicht <u>geerbt</u>	Ez azt jelenti, hogy nem <u>örökölte</u> a problémát, hogy apja és nagyapja láthatóan	<i>Der Patient hilft dem Arzt die Korrekturfunktion zu verwenden; der Arzt korrigiert den Satz mit Hilfe der Tastatur (geärgert - geerbt)</i>

8. Ergebnisse

8.1 BenutzerInnenfreundlichkeit

Über die BenutzerInnenfreundlichkeit von *SayHi Translate* wurden in Kapitel 4.2 bereits einige wichtige Informationen gegeben. Die App ist logisch und transparent aufgebaut und ist daher schon bei der ersten Nutzung einfach und intuitiv zu verwenden. Auch das schlichte und moderne Design in den Farben Blau-Weiß erleichtert die Übersichtlichkeit; die richtige Farbwahl bzw. ein angenehmer Farbkontrast zwischen Hintergrund- und Schriftfarbe kann die BenutzerInnenenerfahrung positiv beeinflussen. Die Anweisungen, die bei der ersten Nutzung der App auf dem Bildschirm erscheinen und die verschiedenen Einstellungen und Funktionen vorstellen, erleichtern ebenfalls die Verwendung. Auch der Patient hat die leichte Handhabung der App im Rahmen des Interviews betont:

Meiner Meinung nach ist die App sehr einfach und logisch aufgebaut. Ich denke, dass ich ein bisschen besser mit Mobilien Apps umgehe als der Durchschnitt meiner Altersklasse, da ich mein Smartphone und Tablet tagtäglich benutze. Deswegen konnte ich die Verwendung der App relativ schnell und einfach lernen, nach ungefähr 15 Minuten Übung habe ich mich gut zurechtgefunden. (IP: 1; übersetzt aus dem Ungarischen von Julia Glocknitzer)

Der Arzt fand den Aufbau der App prinzipiell gut, wies jedoch darauf hin, dass das ständige Hin- und Herdrehen des Tablets während der Untersuchung zeitaufwändig und störend war:

Also im Prinzip ist die App logisch aufgebaut, einfach zu verwenden und gut. Wenn man aber in der Situation ist, dass man miteinander spricht, wäre es natürlich schlau, wenn die App – wie ein Tennisplatz – eine Mitte hat und sie schreibt einmal für mich richtig und einmal für den Gesprächspartner richtig, weil das kostet Zeit, das Hin- und Herdrehen. (IA: 1)

Ich möchte an dieser Stelle anmerken, dass sowohl der Patient als auch der Arzt – trotz der Tatsache, dass beide einer Generation angehören, deren Mitglieder die Verwendung der digitalen Endgeräte erst im Erwachsenenalter erlernt haben – sehr gut mit Smartphones, Tablets und Computern umgehen. Dies hat zum reibungslosen Ablauf des Experiments stark beigetragen. Es ist anzunehmen, dass sich vor allem ältere ProtagonistInnen, die eventuell seltener mit Smartphones in Berührung kommen, mit modernen Technologien und daher auch mit dieser App nicht so leicht zurechtfinden.

- [...] Gibt es vielleicht etwas, was Sie zu der ganzen Situation hinzufügen würden?
- Vielleicht nur, dass ich im Spital...wir haben sehr viele Sachen, die wir diktieren müssen, also wir machen den Operationsbericht, wir machen Arztbriefe, das wird alles über so...elektronische Apps gemacht und ich habe daher gelernt, deutlich zu sprechen, oder sag ma's so... deutlicher zu sprechen, damit der Apparat einen gut verstehen kann. Ich kann mir aber gut vorstellen, dass viele das nicht können und dass dann die Anwendbarkeit der App damit mehr limitiert ist, weil natürlich das Korrigieren, das nicht genaue Verstehen des Apparats sinnverzerrend wirkt und dann ist die Kommunikation gestört. (IA: 7)

Auch die Individualisierbarkeit der Spracheingabe und -ausgabe darf bezüglich der BenutzerInnenfreundlichkeit nicht unerwähnt bleiben. Wie bereits in Kapitel 4.2 dargestellt, gibt es die Möglichkeit, die Sprachausgabe an die eigenen Bedürfnisse anzupassen. Man kann die Geschwindigkeit der Sprachausgabe verändern bzw. auswählen, ob eine männliche oder eine weibliche Stimme bei der Sprachsynthese zu hören ist. Auch die Art und Weise der Spracheingabe kann eingestellt werden. Außerdem können BenutzerInnen bei einigen Sprachen den gewünschten Dialekt auswählen; dies ist vor allem bei Sprachen, die in vielen verschiedenen Ländern gesprochen werden (wie Arabisch, Spanisch und Englisch), relevant und hilfreich.

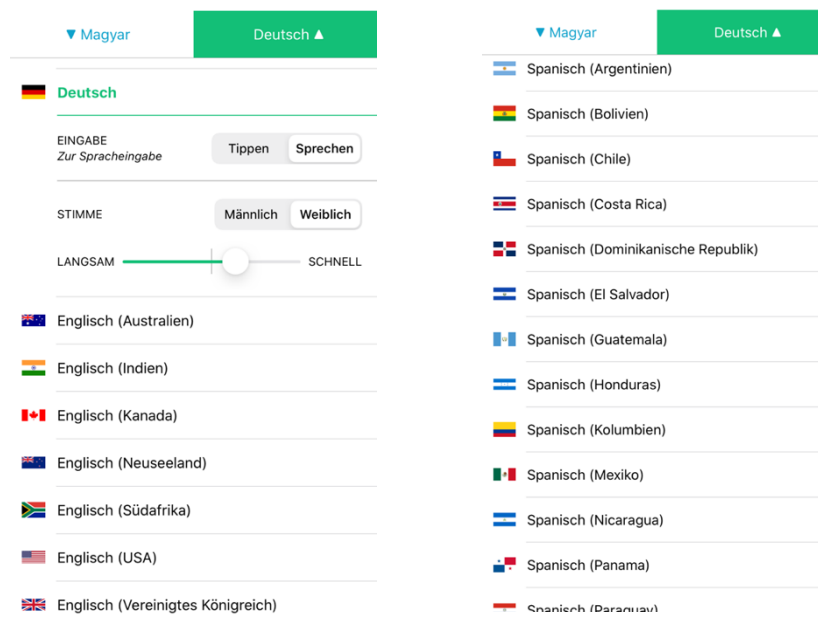


Abbildung 20: Individualisierbarkeit der Spracheingabe bzw. -ausgabe; Vielzahl an Dialekten (Screenshot)

Im Rahmen des Experiments wurden sowohl bei der Spracheingabe als auch bei der Sprachausgabe die Grundeinstellungen beibehalten. Die vom Patienten gesprochenen ungarischen

Sätze wurden von einer männlichen Stimme ins Deutsche „gedolmetscht“, während die Stimme der ungarischen Sprachsynthese weiblich war. Der Patient fand dies jedoch gar nicht störend oder problematisch:

Die Tatsache, dass der Arzt ein Mann war, die App aber eine weibliche Stimme für das Dolmetschen verwendet hat, war gar nicht störend. Das ist auch bei einem „traditionellen“ Dolmetscheinsatz normal. (IP: 2)

Auch die manuelle Korrekturfunktion ist zweifellos eine gute Möglichkeit, das Gesagte zu ergänzen oder die Transkription zu korrigieren, ohne mehrere Sätze wiederholen zu müssen. Wenn man das zu korrigierende Chat-Fenster gedrückt hält und die Option „Bearbeiten“ auswählt, erscheint eine Tastatur, mit deren Hilfe der Text einfach ausgebessert werden kann. Diese Funktion ist vor allem dann hilfreich, wenn aufgrund von Spracherkennungsproblemen das Gesagte von der App nicht richtig transkribiert wurde. Im Laufe des Experiments wurde diese Funktion vom Arzt für die Korrektur falsch erkannter Wörter verwendet, die Erfahrungen diesbezüglich werden im zweiten Evaluierungspunkt, bei den **Spracherkennungsfehlern**, detaillierter dargestellt.

- [...] Haben Sie die manuelle Korrekturfunktion verwendet?
- Ja.
- War sie einfach zu verwenden, oder war es eher störend, immer wieder schreiben zu müssen?
- Es ist dann störend, wenn das Gerät mich nicht gut versteht, aber das Gerät versteht mich meistens nicht gut, wenn ich in dem Inhalt, in dem, was ich sagen möchte, unsicher bin und zaudere. Das merkt dieser Apparat sehr schnell. (IA: 3)

Der Arzt wies mit dieser Aussage auf ein weiteres Problem hin, nämlich darauf, dass die App mehrmals zu früh mit der Aufzeichnung aufhörte und bereits mit dem Dolmetschen anfing, wenn der Redner eine kurze Pause zwischen zwei Wörtern machte. Die automatische Erkennung des Endes einer Spracheingabe ist zwar insgesamt nützlich und kommunikationsbeschleunigend, die Gesprächspartner können sich aber dadurch unter Druck gesetzt fühlen. Es ist leicht zu erkennen, wenn die Transkription vorzeitig beendet wird: Der Beginn und das Ende der Aufnahme werden mit einem Tonsignal gekennzeichnet.

Die Transkriptionen werden in Form von Chat-Fenstern automatisch gespeichert. Dies ist auf jeden Fall positiv zu bewerten und erscheint sehr sinnvoll, da man so das Gespräch später

erneut lesen kann. Da leider schon vor der detaillierten lexikalischen Analyse und Leistungsevaluierung festzustellen war, dass die App ins Ungarische kaum korrekte und verständliche Übersetzungen geliefert hat, war die Speicherung des Gesprächs ausschließlich für die Transkription und Evaluierung nützlich, für den Patienten waren die Übersetzungen auch nach der Untersuchung nicht klar verständlich.

Bezüglich der BenutzerInnenfreundlichkeit ist weiters zu bemerken, dass zur Nutzung der App unbedingt eine stabile Internetverbindung erforderlich ist. Je schwächer das Internet-signal ist, desto langsamer und unzuverlässiger wird die Leistung der App. Da der Patient keine aktive Mobilfunkfunktion mit Internetverbindung auf seinem Tablet hatte, musste er das Gerät mit dem WLAN der Ordination verbinden, um die App verwenden zu können. Im Wartezimmer war das Signal ziemlich schwach, die App brauchte sehr lange, um einen Probesatz zu übersetzen, was in der gegebenen Situation gar nicht beruhigend war. Zum Glück war die Internetverbindung im Ordinationszimmer viel stärker, die Untersuchung konnte also aus Sicht der Verwendbarkeit der App reibungslos durchgeführt werden.

Bezüglich des lautsprachlichen Outputs lässt sich feststellen, dass die Stimme der Sprachsynthese sehr „mechanisch“ wirkt und nicht mit einer echten menschlichen Stimme zu vergleichen ist. Da die Frage- und Aufforderungssätze nicht als solche erkannt wurden, war die Betonung der Wörter sehr monoton und in einigen Fällen sogar falsch. Von der Monotonie und der mechanischen Aussprache abgesehen, war die Sprachausgabe grundsätzlich klar und verständlich.

8.2 Spracherkennungsfehler

Fehler, die bei der automatischen Spracherkennung auftreten, können die Verständlichkeit der Sätze und deren Übersetzungen deutlich beeinträchtigen. Obwohl bei *SayHi Translate* die Möglichkeit gegeben ist, die falsch erkannten Wörter manuell zu korrigieren, werden die Transkriptionen nicht immer von den GesprächspartnerInnen kontrolliert, vor allem weil die Spracheingabe mündlich und nicht mit Hilfe einer Tastatur erfolgt. Auch der Arzt wies nach der Untersuchung darauf hin, dass „das Korrigieren, das nicht genaue Verstehen des Apparats sinnverzerrend wirkt.“ (IA: 7) Obwohl beide Teilnehmer klar und dialektfrei gesprochen haben, hat die App viele Wörter falsch interpretiert. Solche Fehler können leicht zu Sinnverschiebungen zwischen Ausgangs- und Zieltext führen und bedeutende Konsequenzen bezüglich des Informationstransfers und der Verständlichkeit der Übersetzung haben. Natürlich darf man nicht alle Spracherkennungsfehler in einen Topf werfen: Es gibt geringfügige Transkriptionsfehler, die

sich in den Übersetzungen gar nicht widerspiegeln, während schwerwiegende Fehler sogar das Satzverständnis massiv beeinträchtigen können. Es ist also sinnvoll, die Fehler zu gewichten. Im Folgenden werden verschiedene Spracherkennungsfehlertypen, die ich beobachten konnte, mit konkreten Beispielen aus dem Experiment dargestellt. Anfangs werden einige geringfügige Fehler, die kaum oder gar keine Konsequenzen für die Übersetzung hatten, vorgestellt. Danach werden immer schwerere Fehlertypen mit Beispielen aufgelistet, bei denen aufgrund des falsch übersetzten Begriffes die Satzbedeutung nicht mehr nachvollziehbar war.

8.2.1 Geringfügige und mittelschwere Spracherkennungsfehler

8.2.1.1 Zusammenschreibung von zwei Wörtern

Während des Experimentes ist es mehrmals vorgekommen, dass zwei Wörter von der automatischen Spracherkennungskomponente in der Transkription zusammengeschrieben wurden.

Tabelle 6: Transkription des Experimentes, Spracherkennungsfehler (Zeilen 12, 41, 54 und 68)

	Gesprochene Eingabe	ASR + Transkription	MÜ
	A	B	C
12	Elnézést, de ezt nem értettem pontosan. Ismételje meg, kérem <i>pont.</i>	Elnézést, de ezt nemértettem pontosan ismételje meg kérem.	Es tut mir leid, aber ich habe es nicht verstanden, bitte wiederholen Sie es genau, bitte.
41	Durch Bewegung der Muskulatur kommt es zu einem Druck auf die <u>Venen</u> , dieser Druck bringt das Blut zu fließen.	Durchbewegung der Muskulatur kommt es zu einem Druck auf die <u>Wiener</u> dieser Druck bringt das Blut zufließen	Az izmok mozgása van a nyomás a bécsi ez a nyomás hozza a vér áramlását
54	Köszönöm szépen, mindent megértettem. Igyekszem többet mozogni és az életmódomat is úgy alakítani, hogy minél kevesebb rizikó legyen benne.	Köszönömszépen mindent megértettem igyekszem többet mozogni és az életmódomat is úgy alakítani, hogy minél kevesebb Rizikó legyen benne	Vielen Dank Für alles, was ich verstehe, versuche ich, mich mehr zu bewegen und meinen Lebensstil so zu gestalten, dass es so wenig Risiko wie möglich hat.

	Gesprochene Eingabe	ASR + Transkription	MÜ
	A	B	C
68	Nagyon szépen köszönöm a vizsgálatot és a hasznos tanácsot. A gyógyszert be fogom venni, illetve ezt a gyógyszert meg fogom vásárolni és remélem, hogy...	Nagyon szépen köszönöm a vizsgálatot és a hasznos tanácsot a gyógyszert be fogom venni, illetve ezt a gyógyszert megfogom vásárolni és remélem hogy	Vielen Dank für die Untersuchung und hilfreiche Ratschläge, die ich nehmen und kaufen dieses Medikament und hoffen, dass

In Zeile 12 wurde *nem értettem* (= „ich habe es nicht verstanden“) zusammengeschrieben, in der Übersetzung spiegelt sich dieser Schreibfehler jedoch nicht wider.

Aufgrund der Zusammenschreibung von *Durch* und *Bewegung* ist in Zeile 41 in der ungarischen Übersetzung das Wort „durch“ (= „által“) nicht zu finden (*Az izmok mozgása* (41C) bedeutet „Die Bewegung der Muskeln“). Dies führte zu einem kleineren Informationsverlust am Anfang des Satzes.

Auch in Zeile 54 ist der gleiche Fehler mit *Köszönöm szépen* (= „Vielen Dank“) zu beobachten. Der Satz wurde hier zwar falsch übersetzt, Grund dafür ist aber nicht die Zusammenschreibung, sondern das fehlende Komma in der Transkription. *Köszönöm szépen, mindent megértettem.* (A54) bedeutet „Vielen Dank, ich habe alles verstanden.“, stattdessen wurde aber diese Aussage als „Vielen Dank für alles, was ich verstehe“ (C54) übersetzt.

In Zeile 68 ist *meg fogom vásárolni* (= „ich werde es kaufen“) (A68) in der Transkription als *megfogom vásárolni* (B68) zu lesen. In der Übersetzung sind die Verben nicht flektiert, ob dies durch die Zusammenschreibung verursacht wurde, ist unklar. Ich möchte an dieser Stelle noch anmerken, dass *ezt a gyógyszert megfogom* wortwörtlich „ich fasse dieses Medikament an“ bedeutet. Die „willkürliche“ und oft falsche Satzsegmentierung der App kann also auch für unkorrekte Übersetzungen verantwortlich sein. Die fehlenden Satzzeichen stellen ein weiteres Problem der automatischen Spracherkennung dar. Obwohl sowohl der Arzt als auch der Patient relativ oft „Punkt“ und „Komma“ gesagt haben, haben sie doch manchmal vergessen (vor allem bei Fragen), die Satzzeichen mündlich beizufügen.

8.2.1.2 Satzzeichen

Tabelle 7: Transkription des Experimentes, Spracherkennungsfehler (Zeilen 9 und 52)

	Gesprochene Eingabe	ASR + Transkription	MÜ
	A	B	C
9	Haben Sie einen Beruf, bei dem sie viel stehen oder sitzen müssen <i>Komma</i> , tut es im Sommer mehr weh als im Winter?	Haben Sie einen Beruf bei dem sie viel stehen oder sitzen müssen, tut es im Sommer mehr weh als im Winter	Ha van egy munka, ahol meg kell állni vagy ülni sokat, hogy jobban fáj a nyáron, mint télen
52	Haben Sie zu dieser Erklärung irgendwelche Fragen?	Haben Sie zu dieser Erklärung irgendwelche Fragen	Kérdése van ezzel a nyilatkozattal kapcsolatban?

In Zeile 9 wurde von der App nicht erkannt, dass vom Arzt eine Frage gestellt wurde. Die Texteinheit wurde als ein Aussagesatz behandelt, was auch die Wortreihenfolge in der Übersetzung beeinflusst hat. Statt *jobban fáj a nyáron, mint télen (C9)* wäre „*Nyáron jobban fáj, mint télen?*“ die korrekte Übersetzung gewesen. Trotz der Tatsache, dass die ungarische Übersetzung sowohl syntaktisch als auch lexikalisch sehr fehlerhaft ist (wortwörtlich zurückübersetzt: „*Wenn es eine Arbeit gibt, wo man viel stehen bleiben oder sitzen muss, dass es im Sommer mehr weh tut als im Winter*“), hat der Patient die Frage verstanden und beantwortet.

In Zeile 52 wurde zwar in der ungarischen Übersetzung erkannt, dass es sich um eine Frage handelt - was auf jeden Fall positiv zu bewerten ist –, die Wortreihenfolge wurde jedoch nicht ganz angepasst.

8.2.1.3 Auslassungen am Anfang der Spracheingabe

Beginnt man zu früh zu sprechen, kann es vorkommen, dass das System die ersten paar Wörter nicht „mitbekommt“ und den Anfang des Satzes nicht transkribiert. Dies kann manchmal zu merklichem Informationsverlust führen, hier haben jedoch solche Spracherkennungsfehler den Informationstransfer nicht stark beeinflusst.

Tabelle 8: Transkription des Experimentes, Spracherkennungsfehler (Zeilen 27 und 30)

	Gesprochene Eingabe	ASR + Transkription	MÜ
	A	B	C
27	<u>Bitte</u> stellen Sie jetzt das linke Bein mit einem halben Schritt Weg nach vorne weg <i>Punkt</i> .	<u>Wählen</u> Sie jetzt das linke Bein mit einem halben Schritt Weg nach vorne weg.	Most válassza ki a bal lábát egy fél lépéssel el.
30	<u>Stellen Sie bitte</u> jetzt das linke Bein zu mir nach hinten und das rechte einen halben Schritt nach vorne.	Jetzt das linke Bein zu mir nach hinten und das Rechte einen halben Schritt nach vorne	Most a bal láb nekem hátra, és a jobb fél lépés előre

In Zeile 27 kann man sehen, dass das System *Bitte* nicht „gehört“ hat, und die Transkription erst ab „*stellen Sie*“ anfang. Der Patient hat diese Instruktion nicht gut verstanden, Grund dafür war aber ein anderer Spracherkennungsfehler (*stellen* – *wählen*). Der Arzt hat das linke Bein des Patienten selbst einem halben Schritt nach vorne geschoben.

In Zeile 30 wurde zwar *Stellen Sie bitte* (30B) vom System nicht „gehört“, die wichtigen Informationen (das linke Bein nach hinten, das rechte einen halben Schritt nach vorne) wurden jedoch übermittelt. In der Übersetzung fehlt zwar die höfliche Anrede, der Informationstransfer wurde aber nicht beeinflusst.

8.2.1.4 Groß- und Kleinschreibung

Grundsätzlich kann man feststellen, dass *SayHi Translate* sowohl beim Satzanfang als auch bei den einzelnen Wörtern sehr auf die Groß- und Kleinschreibung achtet. Es gab nur wenige Stellen, wo das System die Groß- und Kleinbuchstaben verwechselt hat.

Tabelle 9: Transkription des Experimentes, Spracherkennungsfehler (Zeile 54)

	Gesprochene Eingabe	ASR + Transkription	MÜ
	A	B	C
54	<u>Köszönöm szépen,</u> mindent megértettem. Igyekszem többet mozogni és az életmódomat is úgy alakítani, hogy minél kevesebb rizikó legyen benne.	<u>Köszönömszépen</u> mindent megértettem igyekszem többet mozogni és az életmódomat is úgy alakítani, hogy minél kevesebb Rizikó legyen benne	Vielen Dank Für alles, was ich verstehe, versuche ich, mich mehr zu bewegen und meinen Lebensstil so zu gestalten, dass es so wenig Risiko wie möglich hat.

Aus Sicht der Übersetzung hatte der Schreibfehler in Zeile 54 keine Relevanz, da die Bedeutung des Wortes *rizikó* (= „Risiko“) dadurch nicht geändert wurde und da es in den zwei Sprachen sowieso unterschiedliche grammatische Regeln bezüglich der Groß- und Kleinschreibung gibt. Man könnte denken, dass derartige Fehler aus Sicht der Dolmetschung keine Rolle spielen, da die Übersetzung am Ende durch die Sprachsynthese „vorgelesen“ wird. Man darf aber nicht vergessen, dass im Deutschen das Groß- und Kleinschreiben von einigen Personal- und Possessivpronomina einen wesentlichen Unterschied ausmacht. Ein einziger Buchstabe entscheidet darüber, ob die dritte Person Singular (sie/ihr), die dritte Person Plural (sie/ihr/ihnen) oder die Höflichkeitsform (Sie/Ihr/Ihnen) verwendet wird. Hört man z.B. den Satz „ICH HABE SIE GESEHEN!“, könnte man nicht entscheiden, ob eine Frau, mehrere Menschen oder die angesprochene Person gesehen wurde. Werden von der App nicht die richtigen Pronomina gewählt, kann dies zu Übersetzungsproblemen führen.

Tabelle 10: Transkription des Experimentes, Spracherkennungsfehler (Zeilen 13, 17 und 70)

	Gesprochene Eingabe	ASR + Transkription	MÜ
	A	B	C
13	Ich frage genauer nach, welche Beschwerden Ihnen die Venen in den Beinen machen <i>Punkt</i>. Haben Sie Schmerzen, wenn Sie sich wenig bewegen und viel sitzen <i>Punkt</i>? Haben Sie Schwellungen in den Knöchelregionen der Beine <i>Punkt</i>? Wenn Sie Schwellungen haben Komma, sind es <u>eher</u> auf der Innenseite des Beines oder auf der Außenseite des Beines?	Ich frage genauer nach, welche Beschwerden Ihnen die Venen in den Beinen machen. Haben sie Schmerzen wenn sie sich wenig bewegen und viel sitzen. Haben Sie Schwellungen in den Knöchelregion der Beine. Wenn Sie Schwellungen haben, sind es <u>dir</u> auf der Innenseite des Beines oder auf der Außenseite	Kérem, többet arról, milyen kellemetlen a vénák a lábad neked. Fájdalmai vannak, amikor keveset mozognak és sokat ülnek. Van duzzanat a boka régióban a lábak. Ha van duzzanat, ez akkor a belsejében a láb vagy a külső
17	Bitte zeigen Sie mir an <u>Ihren</u> Beinen die Bereiche, in denen Sie Beschwerden spüren, wenn Sie länger sitzen.	Bitte zeigen Sie mir an <u>ihren</u> Beinen die Bereiche in denen sie Beschwerden spüren, wenn sie länger sitzen	Kérjük, mutassa meg a lábukat a területeken, ahol úgy érezik, kényelmetlenséget, ha ülünk tovább
70	Ich wünsche Ihnen eine gute Heimfahrt, <i>Komma</i> nett, dass Sie bei mir waren <i>Rufzeichen</i>!	Ich wünsche Ihnen eine gute Heimfahrt, nett dass sie bei mir waren!	Szeretném, ha egy jó utazás haza, szép, hogy voltak velem!

In den Zeilen 13, 17 und 70 wurde *Sie* in der Transkription kleingeschrieben, wodurch in der ungarischen Übersetzung die **Verben** nach dritter Person Plural flektiert wurden. Natürlich waren so die Übersetzungen nicht korrekt, zu einer drastischen Sinnverschiebung kam es aber nicht. Der Patient hat in allen drei Fällen den Arzt verstanden und konnte trotz des falschen Pronominalbezugs auf das Gesagte entsprechend reagieren.

Tabelle 11: Transkription des Experimentes, Spracherkennungsfehler (Zeile 60)

	Gesprochene Eingabe	ASR + Transkription	MÜ
	A	B	C
61	Ich habe Sie bei der Untersuchung mehrfach aufgefordert, zu <u>husten</u> <i>Punkt.</i> <u>Husten</u> bedeutet eine Kontraktion mit dem Zwerchfell, dies ist der größte und stärkste Muskel in unserem Körper.	Ich habe sie bei der Untersuchung mehrfach aufgefordert zu <u>husten</u>. <u>Husten</u> bedeutet eine Kontraktion mit dem Zwerchfell, dies ist der größte und stärkste Muskel in unserem Körper	Megkértem őt, hogy köhögést többször a vizsgálat során. Köhögés: összehúzódása a rekeszizom, ez a legnagyobb és legerősebb izom testünk

In Zeile 61 besteht das gleiche Problem, hier bezieht sich jedoch das Personalpronomen *sie* in der Übersetzung auf eine Frau (dritte Person Singular). Auch hier war der Fehler bei der Dolmetschung zwar störend, der Sinn des Satzes blieb aber für den Patienten nachvollziehbar.

Tabelle 12: Transkription des Experimentes, Spracherkennungsfehler (Zeilen 43 und 44)

43	Bei vielen Patienten bestehen <u>Venenprobleme</u>, weil eine Klappe nicht funktioniert <i>Komma</i>, das Blut umdreht und wieder <u>Richtung</u> Beine fließt, wie ich es Ihnen kurz auf der Zeichnung zeigen werde.	Bei vielen Patienten bestehen <u>Bienen Probleme</u>, weil eine Klappe nicht funktioniert, das Blut umdreht und wieder <u>Richtungs</u> Beine fließt wie ich es ihnen kurz auf der Zeichnung zeigen werde	Sok beteg problémái vannak a méhek, mert a fedél nem működik, a vér fordul, és a vér áramlik vissza a lábak felé, mint én megmutatom nekik röviden a rajz
44	Das ist bei Ihnen aber nicht der Fall <i>Komma</i>, die Klappen funktionieren alle <u>sehr gut</u>.	Das ist bei ihnen aber nicht der Fall, die Klappen funktionieren alle <u>sercan</u>	De nem ez a helyzet velük, a füleket minden munkát sercan

Auch in den Zeilen 43 und 44 hat die Kleinschreibung von *Ihnen* zu einem falschen Pronominalbezug (dritte Person Plural) geführt.

8.2.1.5 „Tippfehler“

Unter „Tippfehler“ verstehe ich Spracherkennungsfehler, bei denen das transkribierte Wort nur um ein paar Buchstaben vom ursprünglich gesagten Wort abweicht. Solche Fehler verändern den Inhalt des Textes nicht enorm, führen aber manchmal zur falschen Syntax sowohl in der Transkription als auch in der Übersetzung.

Tabelle 13: Transkription des Experimentes, Spracherkennungsfehler (Zeilen 15 und 18)

	Gesprochene Eingabe	ASR + Transkription	MÜ
	A	B	C
15	Bevor ich Ihnen das <u>beantworte</u>, möchte ich Sie gerne untersuchen <i>Punkt</i>. Dazu werde ich mir <u>Ihre</u> Beine im Stehen anschauen und einen Ultraschall machen <i>Punkt</i>. Dazu müssen Sie bitte in der Garderobe hinter <u>Ihnen</u> die Beine freimachen <i>Punkt</i>. Bitte ziehen Sie Ihre Hose, die Schuhe und die Strümpfe aus und stellen Sie sich auf das kleine <u>Podest</u> <i>Komma</i>, wo die Füße grün eingezeichnet sind.	Bevor ich Ihnen das <u>beantwortet</u> möchte ich Sie gerne untersuchen. Dazu werde ich mir <u>ihre</u> Beine im Stehen anschauen und einen Ultraschall machen. Dazu müssen Sie bitte in der Garderobe hinter <u>ihnen</u> die Beine freimachen. Bitte ziehen Sie Ihre Hose die Schuhe und die Strümpfe aus und stellen Sie sich auf das kleine <u>Protest</u>, wo die Füße grün eingezeichnet sind	Mielőtt válaszolnék nektek, hogy szeretném megvizsgálni téged. Fogom nézni a lábát állva, és nem egy ultrahang. Ehhez kérjük, hogy törölje a lábát az öltözőben mögöttük. Kérjük, vegye le a nadrágját, és harisnya és állni a kis tiltakozás, ahol a lábad készült zöld
18	Ok. Bitte bleiben Sie gleich <u>so stehen</u>, ich werde jetzt an mehreren Stellen <u>eine</u> Ultraschallsonde hinhalten und dabei gleichzeitig auf <u>Ihre</u> Waden drücken <i>Punkt</i>. Es wird Ihnen nichts dabei <u>wehtun</u>.	O. K. Bitte bleiben Sie gleich <u>zustehen</u>, ich werde jetzt an mehreren Stellen <u>in der</u> Ultraschallsonde hinhalten und dabei gleichzeitig auf <u>ihre</u> Waden drücken. Es wird Ihnen nichts dabei <u>wehtuend</u>	Ok. Kérem, maradjon rendben, én most megáll több helyen az ultrahangos szondán, és nyomjuk meg a borjak ugyanabban az időben. Semmi sem fog bántani

In Zeile 15 wurde zur Verbform *beantworte* ein überflüssiger Buchstabe hinzugefügt, wodurch die Transkription syntaktisch inkorrekt wurde, in der Übersetzung ist dieser Fehler aber nicht zu erkennen: das Verb *beantworten* wurde korrekt in erster Person Singular übersetzt (*válaszolnék*).

Der Arzt hat während der Untersuchung stark darauf geachtet, die Wörter einzeln und dialektfrei auszusprechen. Als er aber das Wort „*wehtun*“ artikulierte, hörte es sich ein wenig umgangssprachlich nach „*weh-tuen*“ an, wodurch die App *wehtuend* erkannt hat. Der Satz *Es wird Ihnen nichts dabei wehtun* wurde von App als *Semmi sem fog bántani* übersetzt. Dies bedeutet „*Es wird dich nichts quälen*“, was sich einerseits in der gegebenen Situation komisch anhört, andererseits auf Ungarisch auch ein wenig beunruhigend wirken kann.

8.2.1.6 Partielle Erkennung von Ausdrücken

Es kamen auch Transkriptionsfehler vor, bei denen ein Teil des Originalausdrucks erhalten blieb. Es wurde nur eine Silbe des Wortes falsch verstanden; mehrere Ausdrücke als ein Wort erkannt, oder umgekehrt: ein Ausdruck in mehrere Wörter zerteilt. Die Wirkung solcher Fehler auf die Übersetzung kann man nicht vorhersagen; sie können in einigen Fällen das Satzverständnis merklich beeinträchtigen, es kann jedoch vorkommen, dass die Satzbedeutung – da das Originalwort in irgendeiner Form in der Transkription zu finden ist – nachvollziehbar bleibt.

Tabelle 14: Transkription des Experimentes, Spracherkennungsfehler (Zeilen 11, 18, 32 und 42)

	Gesprochene Eingabe	ASR + Transkription	MÜ
	A	B	C
11	<u>Spüren</u> Sie eine unangenehme... unangenehmen Schmerz in den Beinen wenn sie länger sitzen oder gibt es <u>auch</u> Schwellungen der Beine <i>Punkt</i>? Wenn Sie Schwellungen haben, ist es <u>eher</u> die Innenseite oder die Außenseite?	<u>Spielen</u> Sie eine unangenehme unangenehmen schmerzenden Beinen wenn sie länger sitzen oder gibt es <u>noch</u> Schwellungen der Beine. Wenn Sie Schwellungen haben ist es <u>hier</u> die ine Seite oder die Außenseite	Játsszon kellemetlen kellemetlen lábakkal, amikor hosszabb ideig ülnek, vagy még a lábak duzzanata is fennáll. Ha van duzzanat itt ez a belső oldalán, vagy a külső

	Gesprochene Eingabe	ASR + Transkription	MÜ
	A	B	C
18	Ok. Bitte bleiben Sie gleich so stehen , ich werde jetzt an mehreren Stellen <u>eine</u> Ultraschallsonde hinhalten und dabei gleichzeitig auf <u>Ihre</u> Waden drücken <i>Punkt</i> . Es wird Ihnen nichts dabei <u>wehtun</u> .	O. K. Bitte bleiben Sie gleich zustehen , ich werde jetzt an mehreren Stellen <u>in der</u> Ultraschallsonde hinhalten und dabei gleichzeitig auf <u>ihre</u> Waden drücken. Es wird Ihnen nichts dabei <u>wehtuend</u>	Ok. Kérem, maradjon rendben , én most megáll több helyen az ultrahangos szondán, és nyomjuk meg a borjak ugyanabban az időben. Semmi sem fog bántani

32	Drehen Sie sich jetzt bitte zu mir nach vorne und halten die Hose ein wenig in den Schritt hinein <i>Komma</i> , damit ich bei der Leiste messen kann.	Drehen Sie sich jetzt bitte zu mir nach vorne und halt in die Hose ein wenig in den Schritt hinein, damit ich bei der Leiste messen kann	Most kérem, forduljon hozzám előre, és tartsa a nadrágot egy kicsit az y, hogy tudom mérni a Lágyulás
----	---	---	--

42	Damit das Blut aber nur in die Richtung "Herz" und von der Außenvene zur inneren <u>Vene</u> fließen kann, braucht es eine Venenklappen, die das Blut nur in die richtige Richtung fließen lassen.	Damit das Blut aber nur in die Richtung Herz und von der Außenlinie zur inneren <u>Biene</u> fließen kann, braucht es eine Venenklappen die das Blut nur in die richtige Richtung fließen lassen	De a vér áramlását csak az irányba a szív és a külső vonal a belső méh, akkor szüksége van a véna szelep, amely csak lehetővé teszi a vér áramlását a helyes irányba.
----	---	---	--

In Zeile 11 wurden die Wörter *Schmerz in den* als Partizip 1 des Verbes *schmerzen*, also *schmerzenden*, erkannt. Dies wurde von der App als *kellemetlen* ins Ungarische übersetzt (C11), was „unangenehm“ bedeutet. Der Patient konnte die Frage wegen der sehr schlechten Transkription nicht verstehen und bat den Arzt, das Gesagte zu wiederholen.

In Zeile 18 wurde *bleiben Sie so stehen* als *bleiben Sie zustehend* transkribiert. Die App übersetzte dies als *maradjon rendben* ins Ungarische, was „bleiben Sie in Ordnung“ bedeutet und keinen Sinn ergibt. Trotz dieses Fehlers interpretierte der Patient die Aussage richtig und blieb stehen.

In Zeile 32 wurde *halten die Hose* von der App als *halt in die Hose* erkannt, was sich jedoch in der Übersetzung nicht widerspiegelt: *tartsa a nadrágot* ist tatsächlich die korrekte ungarische Übersetzung des ursprünglichen Ausdruckes. Trotzdem hat der Patient diese Aufforderung nicht gut verstanden, Grund dafür war aber nicht dieser Spracherkennungsfehler, sondern die Übersetzung des Wortes „Schritt“, das einfach als „y“ ins Ungarische übertragen wurde.

Auch die Transkription des Wortes *Außenvene* als *Außenlinie* in Zeile 42 hat den Informationstransfer stark beeinträchtigt.

8.2.2 Schwerwiegende Spracherkennungsfehler

Am Ende dieser Liste befinden sich die schwerwiegenden Spracherkennungsfehler, die zu einer starken Sinnverschiebung zwischen den ursprünglichen Aussagen und ihrer Übersetzungen führten.

Da der Patient auf Krampfaderprobleme hin untersucht wurde, verwendete der Arzt das Wort „Vene“ in unterschiedlichen Varianten und Zusammensetzungen (Venen, Außenvene, Venenblut, Venenklappen usw.) während der Untersuchung insgesamt 15 Mal. Das Wort wurde von der App 7 Mal richtig erkannt und 8 Mal falsch interpretiert. Die falschen Ausdrücke in der Transkription und dadurch auch in der ungarischen Übersetzung führten oft dazu, dass der Patient den Faden verlor.

Tabelle 15: Transkription des Experimentes, Spracherkennungsfehler (Zeilen 37, 38, 41 und 59)

	Gesprochene Eingabe	ASR + Transkription	MÜ
	A	B	C
37	Es gibt eine große Vene ganz in der Tiefe der Muskulatur, die das meiste Blut sammelt und direkt zum Herzen führt <i>Punkt</i> .	Gibt eine große Biene ganz in der Tiefe der Muskulatur, die das meiste Blut sammelt und direkt zum Herzen führen.	Ad egy nagy méh a mélység az izmok, hogy összegyűjti a legtöbb vért, és vezet közvetlenül a szívbe.
38	Diese tiefe Beinvene wird unterstützt durch zwei dünne Venen , die unter der Haut laufen und das Blut von der <u>Hautregion</u> einsammeln.	Diese tiefe Beinvene wird unterstützt durch zwei dünne Bienen die unter der Haut laufen und das Blut von der <u>Haut Region</u> einsammeln	Ez a mély láb véna támasztja alá két vékony méhek fut a bőr alatt, és összegyűjti a vért a bőr régió

	Gesprochene Eingabe	ASR + Transkription	MÜ
	A	B	C
41	<u>Durch Bewegung</u> der Muskulatur kommt es zu einem Druck auf die <u>Venen</u> , dieser Druck bringt das Blut zu fließen.	<u>Durchbewegung</u> der Muskulatur kommt es zu einem Druck auf die <u>Wiener</u> dieser Druck bringt das Blut zufließen	Az izmok mozgása van a nyomás a <u>bécsi</u> ez a nyomás hozza a vér áramlását

59	Wenn Sie ins Fitnessstudio gehen, Sport betreiben oder sich körperlich anstrengen, kommt jetzt noch von mir ein wichtiger Tipp <i>Komma</i> , um ihre <u>Venen</u> auch in Zukunft gesund zu halten.	Wenn Sie ins Fitnessstudio gehen, Sport betreiben oder sich körperlich anstrengen kommt jetzt noch von mir ein wichtiger Tipp, um ihre <u>Bienen</u> auch in Zukunft gesund zu halten	Ha megy a tornaterem, a testmozgás és a testmozgás kemény, adok egy fontos tipp, hogy tartsa a <u>méhek</u> egészségesek a jövőben
----	--	---	--

Wie man Anhand der Beispiele sehen kann, wurde vom Spracherkennungssystem statt *Venen* oft *Bienen* verstanden (6 Mal). Der Patient hatte wahrscheinlich keine Ahnung, warum während einer ärztlichen Untersuchung so oft über Bienen („méhek“) gesprochen wurde, fragte aber nicht nach.

[...], wegen der Fehlübersetzungen habe ich aber leider mehrmals den Faden verloren, da das Gerät nicht einmal annähernd das übersetzt hat, worüber wir gesprochen haben. Hätte ich ernstere Beschwerden gehabt, hätte ich wahrscheinlich nicht verstanden, was der Arzt mir mitteilen wollte. Oder vielleicht hätte er mich nicht verstanden. (IP: 4)

In Zeile 41 wurde *Venen* als *Wiener* transkribiert und als *bécsi* ins Ungarische übersetzt. Auch dieser Fehler führte folglich zu einer massiven Sinnverschiebung.

Tabelle 16: Transkription des Experimentes, Spracherkennungsfehler (Zeile 15)

	Gesprochene Eingabe	ASR + Transkription	MÜ
	A	B	C
15	Bevor ich Ihnen das <u>beantworte</u>, möchte ich Sie gerne untersuchen <i>Punkt</i>. Dazu werde ich mir <u>Ihre</u> Beine im Stehen anschauen und einen Ultraschall machen <i>Punkt</i>. Dazu müssen Sie bitte in der Garderobe hinter <u>Ihnen</u> die Beine freimachen <i>Punkt</i>. Bitte ziehen Sie Ihre Hose, die Schuhe und die Strümpfe aus und stellen Sie sich auf das kleine <u>Podest</u> <i>Komma</i>, wo die Füße grün eingezeichnet sind.	Bevor ich Ihnen das <u>beantwortet</u> möchte ich Sie gerne untersuchen. Dazu werde ich mir <u>ihre</u> Beine im Stehen anschauen und einen Ultraschall machen. Dazu müssen Sie bitte in der Garderobe hinter <u>ihnen</u> die Beine freimachen. Bitte ziehen Sie Ihre Hose die Schuhe und die Strümpfe aus und stellen Sie sich auf das kleine <u>Protest</u>, wo die Füße grün eingezeichnet sind	Mielőtt válaszolnék nektek, hogy szeretném megvizsgálni téged. Fogom nézni a lábát állva, és nem egy ultrahang. Ehhez kérjük, hogy törölje a lábát az öltözőben mögöttük. Kérjük, vegye le a nadrágját, és harisnya és állni a kis <u>tiltakozás</u>, ahol a lábad készült zöld

Auch die falsche Erkennung des Wortes *Podest* als *Protest* veränderte die Satzbedeutung in Zeile 15 sehr und ergab einen unverständlichen Satz. Der Patient verstand zwar, dass eine Untersuchung folgen würde, für die er in der Garderobe die Hose ausziehen sollte; wohin er sich aber nachher stellen sollte, blieb ihm unklar. Der Arzt löste das Problem, indem er auf das Podest zeigte.

Tabelle 17: Transkription des Experimentes, Spracherkennungsfehler (Zeilen 44 und 45)

	Gesprochene Eingabe	ASR + Transkription	MÜ
	A	B	C
44	Das ist bei <u>Ihnen</u> aber nicht der Fall <i>Komma</i>, die Klappen funktionieren alle <u>sehr gut</u>.	Das ist bei <u>ihnen</u> aber nicht der Fall, die Klappen funktionieren alle <u>sercan</u>	De nem ez a helyzet velük, a füleket minden munkát <u>sercan</u>
45	Die Klappen funktionieren bei Ihnen alle <u>sehr gut</u>	Die Klappen funktionieren bei Ihnen alle sehr gut	A füleket <u>nagyon jól</u> működik az Ön számára

In Zeile 44 verstand die App *sehr gut* nicht richtig und erkannte es als *sercan*, was in keiner der beiden Sprachen einen Sinn ergibt. Der Arzt bemerkte den Fehler und wiederholte den Satz noch einmal, wobei er dann *sehr gut* klarer und lauter aussprach (45A).

Tabelle 18: Transkription des Experimentes, Spracherkennungsfehler (Zeilen 46 – 49)

	Gesprochene Eingabe	ASR + Transkription	MÜ	Bildbeschreibung, nonverbale Kommunikation
	A	B	C	D
46	Damit haben Sie das Problem <i>Komma</i> , dass ihr Vater und ihr Großvater offenbar gehabt haben, nicht geerbt .	Damit haben Sie das Problem, dass ihr Vater und ihr Großvater offenbar gehabt haben nicht geärgert	Ez azt jelenti, hogy a probléma, hogy apja és nagyapja láthatóan nem bosszantott	
47	Nicht geerbt .	Nicht geärgert	Nem bosszús	<i>erkennt, dass das Ende des Satzes nicht richtig transkribiert wurde und wiederholt die letzten zwei Wörter</i>
48				<i>zeigt "es spinnt"</i>
49		Damit haben Sie das Problem, dass ihr Vater und ihr Großvater offenbar gehabt haben nicht geerbt	Ez azt jelenti, hogy nem örökölte a problémát, hogy apja és nagyapja láthatóan	<i>Der Patient hilft dem Arzt die Korrekturfunktion zu verwenden; der Arzt korrigiert den Satz mit Hilfe der Tastatur (geärgert - geerbt)</i>

Auch mit der Transkription des Wortes *geerbt* gab es Probleme: In Zeile 46 wurde das Verb falsch erkannt und als *geärgert* transkribiert. Der Arzt bemerkte den Fehler und wiederholte die letzten zwei Wörter des Satzes in einer neuen Spracheingabe (Zeile 47). Da aber auch so *geärgert* von der App verstanden wurde, musste der Patient eingreifen und zeigen, wie die manuelle Korrekturfunktion verwendet wird. Mit Hilfe der Tastatur konnte der Arzt das Wort in der Transkription korrigieren, und der Patient bekam endlich das beruhigende Ergebnis: Er hat die Krankheit seines Vaters und Großvaters glücklicherweise nicht *geerbt*.

Tabelle 19: Transkription des Experimentes, Spracherkennungsfehler (Zeilen 19 – 26)

	Gesprochene Eingabe	ASR + Transkription	MÜ	Bildbeschreibung, nonverbale Kommunikation
	A	B	C	D
19	Bitte atmen Sie tief ein und husten Sie, wenn ich jetzt sage.	Bitte atmen Sie tief ein und wussten Sie, wenn ich jetzt sage	Kérjük, hogy egy mély lélegzetet, és tudom , mikor azt mondom, most	<i>Patient atmet tief ein</i>
20	Jetzt!	x	x	
21				<i>atmet tief ein</i>
22				<i>hustet, um zu zeigen, was er gemeint hat</i>
23				<i>hustet nach</i>
24	Ok, Stopp.	x	x	
25	Jetzt!	x	x	
26				<i>hustet</i>

Auch das Wort „husten“ stellte für das Spracherkennungssystem eine Herausforderung dar. Die falsche Erkennung und Übersetzung des Wortes führte in Zeile 19 zu einem Kommunikationsproblem, das gelöst werden musste. Da *husten* als *wussten* erkannt wurde, gab es in der Übersetzung einen bedeutenden Informationsverlust. Da nur die Aufforderung *atmen Sie tief ein* korrekt übersetzt wurde, atmete der Patient beim *Jetzt!* (A20) automatisch tief ein, hustete aber nicht. Der Arzt löste das Problem, indem er selbst hustete, um zu zeigen, was gemeint war. Danach wusste der Patient Bescheid, was zu tun war.

Tabelle 20: Transkription des Experimentes, Spracherkennungsfehler (Zeilen 60 – 63)

	Gesprochene Eingabe	ASR + Transkription	MÜ	Bildbeschreibung, nonverbale Kommunikation
	A	B	C	D
60	Ich habe <u>Sie</u> bei der Untersuchung mehrfach aufgefordert, zu <u>husten</u> Punkt. <u>Husten</u> bedeutet eine Kontraktion mit dem Zwerchfell, dies ist der größte und stärkste Muskel in unserem Körper.	Ich habe <u>sie</u> bei der Untersuchung mehrfach aufgefordert zu <u>wussten</u> . <u>Houston</u> bedeutet eine Kontraktion mit dem Zwerchfell, dies ist der größte und stärkste Muskel in unserem Körper	Kérdeztem tőle több idő alatt a vizsgálat-hoz <u>tud</u> . <u>Houston</u> azt jelenti, összehúzódása a rekeszizom, ez a legnagyobb és legerősebb izom testünk.	
61	Ich habe <u>Sie</u> bei der Untersuchung mehrfach aufgefordert, zu <u>husten</u> Punkt. <u>Husten</u> bedeutet eine Kontraktion mit dem Zwerchfell, dies ist der größte und stärkste Muskel in unserem Körper.	Ich habe <u>sie</u> bei der Untersuchung mehrfach aufgefordert zu <u>husten</u> . <u>Husten</u> bedeutet eine Kontraktion mit dem Zwerchfell, dies ist der größte und stärkste Muskel in unserem Körper	Megkértem őt, hogy <u>köhögést</u> többször a vizsgálat során. <u>Köhögés</u> : összehúzódása a rekeszizom, ez a legnagyobb és legerősebb izom testünk	<i>verwendet die manuelle Korrekturfunktion, um die falsch erkannten Wörter (wussten - husten, Houston - Husten) zu korrigieren</i>
62	Na!	x	x	<i>ärgert sich</i>
63	Na, jetzt geht's.	x	x	

Dasselbe Problem kam auch in Zeile 60 vor, in der *husten* zuerst als *wussten* und dann als *Houston* erkannt wurde. Da der Arzt die manuelle Korrekturfunktion schon bedienen konnte, konnte er diese Fehler schnell korrigieren.

Tabelle 21: Transkription des Experimentes, Spracherkennungsfehler (Zeilen 66 – 67)

	Gesprochene Eingabe	ASR + Transkription	MÜ
	A	B	C
66	Damit die Venen nicht krank werden <i>Komma</i>, achten Sie bitte bei sportlicher Belastung und körperlicher Anstrengung immer darauf, dass sie nicht mit dem Bauch pressen <i>Punkt</i>. Dies tut man üblicherweise, um sich die Belastung zu erleichtern <i>Punkt</i>. Es macht aber einen Druck auf die Venenklappen <i>Komma</i>, und kann eventuell auch noch später zu Problemen mit den <u>Venen</u> führen <i>Punkt</i>.	Damit die Venen nicht krank werden, achten Sie bitte bei sportlicher Belastung und körperliche Anstrengung immer darauf, dass sie nicht mit dem Bauchpressen. Dies tut man üblicherweise, um sich die Belastung zu erleichtern. Es macht aber einen Druck auf die Venenklappen, und kann eventuell auch noch später zu Problemen mit den <u>Bienen</u> führen.	Annak megakadályozására, hogy a vénák szerzés beteg, kérjük, mindig győződjön meg arról, hogy nem nyomja a has, ha gyakorlása és a kifejtése. Ez általában történik, hogy enyhítse a törzs. Ez azonban nyomást gyakorol a véna flaps, és problémákat okozhat a <u>méhek</u> még később.
67		Damit die Venen nicht krank werden, achten Sie bitte bei sportlicher Belastung und körperliche Anstrengung immer darauf, dass sie nicht mit dem Bauchpressen. Dies tut man üblicherweise, um sich die Belastung zu erleichtern. Es macht aber einen Druck auf die Venenklappen, und kann eventuell auch noch später zu Problemen mit den <u>Venen</u> führen.	Annak megakadályozására, hogy a vénák szerzés beteg, kérjük, mindig győződjön meg arról, hogy nem nyomja a has, ha gyakorlása és a kifejtése. Ez általában történik, hogy enyhítse a törzs. Ez azonban nyomást gyakorol a véna flaps, és problémákat okozhat a <u>vénák</u> még később.

Wie vorher erläutert, hatte die App mit der Erkennung des Wortes *Venen* sehr oft Probleme. Dieser Fehler wurde jedoch nur einmal, in den letzten Minuten der Untersuchung, vom Arzt wahrgenommen (Zeile 66). Dies war das dritte Mal, dass er die Transkription manuell korrigierte.

8.3 Übersetzungsfehler

Wie im vorigen Kapitel dargestellt, war die falsche Transkription des Gesagten oft der Grund, warum die App fehlerhafte, inkorrekte und manchmal sogar unverständliche Übersetzungen lieferte. In diesem Kapitel möchte ich mich auf jene Fehler beschränken, die nicht aufgrund der falschen Spracherkennung, sondern während des maschinellen Übersetzungsprozesses entstanden. Auch hier erscheint eine Gewichtung der Fehler sinnvoll. Dieser Teil stützt sich auf das Gewichtungsschema von Cap (2003: 61), in dem die Übersetzungsfehler in vier unterschiedliche Kategorien eingeteilt werden¹³:

Tabelle 22: Beispiel für ein Gewichtungsschema nach Cap (2003: 61)

Gewichtung	Beschreibung
1	Geringfügiger Fehler, der keinerlei Auswirkung auf die Bedeutung des Satzes hat (zum Beispiel Orthographiefehler, falscher Artikel, stilistischer Fehler)
2	Fehler, der die Bedeutung des Satzes nicht wesentlich beeinträchtigt (zum Beispiel falsche Präposition, falsche Wortstellung)
3	Fehler, der das Satzverständnis merklich beeinträchtigt (zum Beispiel Pronominalbezug)
4	Schwerwiegender Fehler, der das Satzverständnis so massiv beeinträchtigt, dass die Satzbedeutung nicht mehr nachvollziehbar ist (zum Beispiel falsche Wortwahl, falscher Ausdruck)

8.3.1 Geringfügige Übersetzungsfehler

8.3.1.1 Orthographiefehler

Es ist auf jeden Fall positiv zu bewerten, dass *SayHi Translate* sowohl auf die Groß- und Kleinschreibung als auch auf die Rechtschreibung der Wörter in den Übersetzungen achtet. Es kamen kaum derartige Fehler vor.

¹³ Der Grund, warum dieses konkrete Gewichtungsschema nicht auch bei den Spracherkennungsfehlern verwendet wurde, ist, dass dort die verschiedenen Fehlerkategorien nicht immer eindeutig gewichtet werden konnten: z.B. waren die Fehler bezüglich der Groß- und Kleinschreibung der Wörter aus Sicht der automatischen Spracherkennung eher geringfügig, sie haben jedoch in einigen Fällen das Satzverständnis in den Übersetzungen merklich beeinträchtigt (z.B. sie – Sie) (vgl. Kapitel 8.2.1.4).

Tabelle 23: Transkription des Experimentes, Übersetzungsfehler, Orthographiefehler (Zeile 64)

	Gesprochene Eingabe	ASR + Transkription	MÜ
	A	B	C
64	Beim Husten oder <u>Pressen</u> mit dem Bauch wird mit großem Druck das Venenblut wieder in die Beine zurückgedrückt.	Beim Husten oder <u>pressend</u> mit dem Bauch wird mit großem Druck das Venenblut wieder in die Beine zurückgedrückt	Amikor köhöt vagy présel a hasüregbe, a véna vére nagy nyomással visszateredik a lábakra

In Zeile 64 wurde statt *köhög* (= „husten“) *köhöt* (64C) geschrieben und von der App „vorgelesen“. Dieser kleine Rechtschreibungsfehler hat jedoch das Wortverständnis nicht beeinflusst.

8.3.1.2 Stilistische Fehler

Tabelle 24: Transkription des Experimentes, Übersetzungsfehler, Stilistische Fehler (Zeilen 12 und 16)

	Gesprochene Eingabe	ASR + Transkription	MÜ
	A	B	C
12	Elnézést, de ezt <u>nem értettem</u> pontosan. Ismételje meg, kérem pont.	Elnézést, de ezt <u>nemértettem</u> pontosan ismételje meg kérem .	Es tut mir leid, aber ich habe es nicht verstanden, bitte wiederholen Sie es genau, bitte .
16	Jó, rendben.	Jó rendben	Alles in Ordnung, in Ordnung.

Die deutsche Übersetzung in Zeile 12 (12C) scheint ein wenig zu höflich zu sein, da das Wort *kérem* (= „bitte“) zweimal übersetzt wurde, die Satzbedeutung wird aber dadurch gar nicht beeinträchtigt. Obwohl es das Satzverständnis nicht stark verändert, möchte ich an dieser Stelle erwähnen, dass sich das Wort „genau“ am falschen Platz befindet, die korrekte Übersetzung wäre „*Es tut mir leid, aber ich habe es nicht genau verstanden. Bitte wiederholen Sie es.*“. Dieser Übersetzungsfehler entstand – wie viele andere – durch das fehlende Komma und die falsche automatische Satzsegmentierung der App.

In Zeile 16 werden sowohl *jó* (= „gut“) als auch *rendben* (= „in Ordnung“) als *in Ordnung* übersetzt (16C), diese Wortwiederholung hört sich zwar komisch an, ändert aber nicht die Bedeutung des Satzes.

Tabelle 25: Transkription des Experimentes, Übersetzungsfehler, Stilistische Fehler (Zeile 61)

	Gesprochene Eingabe	ASR + Transkription	MÜ
	A	B	C
61	Ich habe <u>Sie</u> bei der Untersuchung mehrfach aufgefordert, zu <u>husten</u> Punkt. <u>Husten bedeutet</u> eine Kontraktion mit dem Zwerchfell, dies ist der größte und stärkste Muskel in unserem Körper.	Ich habe <u>sie</u> bei der Untersuchung mehrfach aufgefordert zu <u>husten</u>. <u>Husten</u> bedeutet eine Kontraktion mit dem Zwerchfell, dies ist der größte und stärkste Muskel in unserem Körper	Megkértem őt, hogy köhögést többször a vizsgálat során. <u>Köhögés:</u> összehúzódása a rekeszizom, ez a legnagyobb és legerősebb izom testünk

In Zeile 61 kann man sehen, dass das deutsche Wort *bedeutet* einfach mit einem Doppelpunkt ins Ungarische übersetzt wurde („*Husten bedeutet*“ – „*Köhögés:*“). Dies kann zwar eine Übersetzungsmöglichkeit sein, wenn es sich um schriftliche Texte handelt, beim Dolmetschen ist diese Lösung aber äußerst ungewöhnlich und klingt auch unnatürlich.

8.3.2 Fehler, die die Bedeutung des Satzes nicht wesentlich beeinträchtigen

8.3.2.1 Falsche Präpositionen

Tabelle 26: Transkription des Experimentes, Übersetzungsfehler, Falsche Präpositionen (Zeilen 30 und 34)

	Gesprochene Eingabe	ASR + Transkription	MÜ
	A	B	C
30	Stellen Sie bitte jetzt das linke Bein <u>zu mir</u> nach hinten und das rechte einen halben Schritt nach vorne.	Jetzt das linke Bein <u>zu mir</u> nach hinten und das Rechte einen halben Schritt nach vorne	Most a bal láb <u>nekem</u> hátra, és a jobb fél lépés előre
34	Danke, mit der Untersuchung sind wir fertig <i>Komma</i>, Sie können sich wieder anziehen und setzen sich <u>zu mir an den Schreibtisch</u>.	Danke mit der Untersuchung sind wir fertig, Sie können sich wieder anziehen und setzen sich <u>zu mir an den Schreibtisch</u>	Köszönöm a vizsgálatot végeztünk, akkor újra felöltözni, és üljön le <u>velem</u> az <u>írásztalon</u>

In Zeile 30 wurde *zu mir nach hinten* statt „*hózzám hátra*“ als *nekem hátra* (30C) übersetzt, was „*mir nach hinten*“ bedeutet und wahrscheinlich erneut durch die falsche Segmentierung des Satzes verursacht wurde.

Auch in Zeile 34 stellte *zu mir* eine Herausforderung für die MÜ dar und wurde als *velem* (= „*mit mir*“) ins Ungarische übertragen. Grundsätzlich wäre auch „*setzen Sie sich mit mir an den Schreibtisch*“ gut verständlich, ein weiterer Übersetzungsfehler erschwerte aber die Verständlichkeit des Satzes. *An den Schreibtisch* wurde als *íróasztalon* (= „*auf dem Schreibtisch*“) übersetzt, in der Übersetzung ist also wortwörtlich „*Setzen Sie sich mit mir auf dem Schreibtisch*“ zu lesen. Der Patient verstand den Satz trotz der Fehler und setzte sich nach der Untersuchung zurück an den Schreibtisch.

8.3.2.2 Deklinationsfehler

Deklination bezieht sich auf die Art und Weise, wie bestimmte Wörter gemäß den grammatischen Kategorien (Kasus, Numerus, Genus) ihre Form verändern. Eine falsche Deklination verändert den Sinn eines Satzes nicht enorm, kann aber zu Ungenauigkeiten in der Übersetzung führen.

Tabelle 27: Transkription des Experimentes, Übersetzungsfehler, Deklinationsfehler (Zeilen 7 und 30)

	Gesprochene Eingabe	ASR + Transkription	MÜ
	A	B	C
7	Darf ich fragen, ist es das rechte Bein oder das linke Bein <i>Punkt</i> , und haben Sie schon irgendeine Form von Behandlung deswegen gehabt?	Darf ich fragen, ist es das rechte Bein oder das linke Bein . Und haben Sie schon irgendeine Form von Behandlung deswegen gehabt	Megkérdezhetem, ez a jobb láb vagy a bal lábát . És volt már bármilyen formában a kezelés e
30	Stellen Sie bitte jetzt das linke Bein zu mir nach hinten und das rechte einen halben Schritt nach vorne.	Jetzt das linke Bein zu mir nach hinten und das Rechte einen halben Schritt nach vorne	Most a bal láb nekem hátra, és a jobb fél lépés előre

Sowohl in Zeile 7 als auch in Zeile 30 wurde *das linke Bein* in den falschen Kasus übertragen. In Zeile 7 (7B) ist *das linke Bein* im Nominativ, wurde aber als *a bal lábát* (Dativ) übersetzt (7C). In Zeile 30 wurde der gleiche Übersetzungsfehler umgekehrt gemacht: *das linke Bein* im Dativ (30B) wurde als *a bal láb* (Nominativ) ins Ungarische übertragen. Ohne den Zusammenhang zu kennen, könnte man natürlich nicht entscheiden, ob *das linke Bein* im Nominativ oder im Dativ gemeint ist, die MÜ-Komponente der App sollte aber aufgrund des Kontextes entscheiden können, welcher Kasus in der Übersetzung zu verwenden ist. Möglicherweise ist der Grund der falschen Deklination, dass in Zeile 30 *Stellen Sie bitte* (30A) nicht transkribiert wurde. Hier wurde auch *das rechte (Bein)* im Dativ (30A) statt a „jobb“ als *a jobb* (Nominativ) übersetzt (30C).

8.3.2.3 Falsche Wortstellung

Tabelle 28: Transkription des Experimentes, Übersetzungsfehler, Falsche Wortstellung (Zeilen 4 und 39)

	Gesprochene Eingabe	ASR + Transkription	MÜ
	A	B	C
4	Gyorsan elmagyarázom, hogy hogyan működik a program. Itt vannak a nyelvek, amikor beszélünk, ezeket a gombokat kell megnyomni és utána kell beszélni pont.	Gyorsan elmagyarázom hogy hogyan működik a program, itt vannak a nyelvek, amikor beszélünk ezeket a gombokat kell megnyomni és utána kell beszélni.	Ich erkläre schnell, wie das Programm funktioniert, hier sind die Sprachen, wenn wir über diese Tasten sprechen, die gedrückt werden und dann sprechen.
39	Es ist nicht selbstverständlich, dass das Blut gegen die Schwerkraft hinauf zum Herzen fließt.	Es ist nicht selbstverständlich dass das Blut gegen die Schwerkraft hinauf zum Herzen fließt	Ez nem magától értetődő, hogy a vér áramlik fel a szívbe a gravitáció ellen

Die falsche Wortstellung ist ein Fehler, der in den meisten von der App gelieferten Übersetzungen mehr oder weniger zu beobachten ist. In Zeile 4 ist *hier sind die Sprachen, wenn wir über diese Tasten sprechen, die gedrückt werden und dann sprechen* (4C) zu lesen. Die korrekte

Übersetzung wäre jedoch „*Hier sind die Sprachen, wenn wir sprechen, müssen wir diese Tasten drücken und erst danach sprechen*“. Die Übersetzung blieb trotz des Fehlers verständlich, der Arzt konnte die Applikation gut bedienen.

Auch in den ungarischen Übersetzungen kommt sehr oft ein falscher Satzbau vor, zum Beispiel in Zeile 39, wo *hogy a vér áramlik fel a szívbe a gravitáció ellen* zu lesen ist, die korrekte Wortstellung aber „*hogy a vér a gravitáció ellen felfelé áramlik a szívbe*“ lauten müsste. Wortstellungsfehler verändern die Satzbedeutung zwar nicht massiv, können aber die Verständlichkeit und die Flüssigkeit der Übersetzung beeinträchtigen.

8.3.3 Fehler, die das Satzverständnis merklich beeinträchtigen

8.3.3.1 Pronominalbezugsfehler

Tabelle 29: Transkription des Experimentes, Übersetzungsfehler, Pronominalbezugsfehler (Zeilen 5 und 13)

	Gesprochene Eingabe	ASR + Transkription	MÜ
	A	B	C
5	Gut, ich glaube, das habe ich verstanden <i>Punkt</i> . Darf ich Sie fragen, was ich für Sie tun kann?	Ich glaube das habe ich verstanden. Darf ich Sie fragen was ich für Sie tun kann	Azt hiszem, megértettem. Megkérdezhetem, hogy mit tehetek érted
13	Ich frage genauer nach, welche Beschwerden Ihnen die Venen in den Beinen machen <i>Punkt</i> . Haben Sie Schmerzen, wenn <u>Sie</u> sich wenig bewegen und viel sitzen <i>Punkt</i> ? Haben Sie Schwellungen in den Knöchelregionen der Beine <i>Punkt</i> ? Wenn Sie Schwellungen haben <i>Komma</i> , sind es <u>eher</u> auf der Innenseite des Beines oder auf der Außenseite <u>des Beines</u> ?	Ich frage genauer nach, welche Beschwerden Ihnen die Venen in den Beinen machen. Haben sie Schmerzen wenn <u>sie</u> sich wenig bewegen und viel sitzen. Haben Sie Schwellungen in den Knöchelregion der Beine. Wenn Sie Schwellungen haben, sind es <u>dir</u> auf der Innenseite des Beines oder auf der Außenseite	Kérem, többit arról, milyen kellemetlen a vénák a lábad neked . Fájdalma van, amikor keveset mozognak és sokat ülnek. Van duzzanat a boka régióban a láb. Ha van duzzanat, ez akkor a belsejében a láb vagy a külső

Obwohl die automatische Spracherkennung in Zeile 5 die höfliche Anrede *Sie* erkannt hat und diese korrekt transkribiert wurde, wird der Patient in der Übersetzung in der Du-Form angesprochen (*mit tehetek érted* – „*Was kann ich für dich tun*“). Dasselbe Problem entstand auch in

Zeile 13, wo *Ihnen* als *neked* (= „dir“) übersetzt wurde. Es kam während der Untersuchung mehrmals vor, dass trotz der richtigen Erkennung der Höflichkeitsform in der Übersetzung die Du-Form verwendet wurde. Dies führt prinzipiell nicht zu einer Sinnverschiebung, kann jedoch als unhöflich interpretiert werden.

Tabelle 30: Transkription des Experimentes, Übersetzungsfehler, Pronominalbezugsfehler (Zeilen 15 und 19)

	Gesprochene Eingabe	ASR + Transkription	MÜ
	A	B	C
15	Bevor ich Ihnen das <u>beantworte</u>, möchte ich Sie gerne untersuchen <i>Punkt</i>. Dazu werde ich mir ihre Beine im Stehen anschauen und einen Ultraschall machen <i>Punkt</i>. Dazu müssen Sie bitte in der Garderobe hinter Ihnen die Beine freimachen <i>Punkt</i>. Bitte ziehen Sie Ihre Hose, die Schuhe und die Strümpfe aus und stellen Sie sich auf das kleine <u>Podest</u> <i>Komma</i>, wo die Füße grün eingezeichnet sind.	Bevor ich Ihnen das <u>beantwortet</u> möchte ich Sie gerne untersuchen. Dazu werde ich mir ihre Beine im Stehen anschauen und einen Ultraschall machen. Dazu müssen Sie bitte in der Garderobe hinter ihnen die Beine freimachen. Bitte ziehen Sie Ihre Hose die Schuhe und die Strümpfe aus und stellen Sie sich auf das kleine <u>Protest</u>, wo die Füße grün eingezeichnet sind	Mielőtt válaszolnék nektek, hogy szeretném megvizsgálni téged. Fogom nézni a lábát állva, és nem egy ultrahang. Ehhez kérjük, hogy törölje a lábát az öltözőben mögöttük. Kérjük, vegye le a nadrágját, és harisnya és állni a kis tiltakozás, ahol a lábad készült zöld
19	Bitte atmen Sie tief ein und <u>husten</u> Sie, wenn ich jetzt sage.	Bitte atmen Sie tief ein und <u>wussten</u> Sie, wenn ich jetzt sage	Kérjük, hogy egy mély lélegzetet, és tudom, mikor azt mondom, most

Auch in Zeile 15 wird der Patient geduzt, *ich möchte Sie gerne untersuchen* wurde statt „*szeretném megvizsgálni Önt*“ als *szeretném megvizsgálni téged* (= „*ich möchte dich untersuchen*“) übersetzt. Würde der Arzt den Patienten tatsächlich per Du ansprechen, würde die grammatisch korrekte Übersetzung „*szeretnélek megvizsgálni téged*“ lauten, die richtige Transkription der Höflichkeitsform *Sie* hat sich also nur in der Konjugation des Verbes widerspiegelt, wurde jedoch als *téged* (= „*dich*“) ins Ungarische übertragen. In derselben Zeile wurde *bevor*

ich **Ihnen** das beantworte als *Mielőtt válaszolnék nektek* (= „Bevor ich euch das beantworte“) übersetzt, was das Satzverständnis weiter erschwert.

Sowohl in Zeile 15 als auch in Zeile 19 stellt das Wort *Bitte* eine Herausforderung für die Übersetzungskomponente dar. *Bitte* wurde statt *Kérem* (= „ich bitte Sie“) als *Kérjük* (= „wir bitten Sie“) übersetzt, was zwar grundsätzlich das Satzverständnis nicht beeinträchtigt, trotzdem ein wenig verwirrend wirken kann. Das Wort *Bitte* wurde während des Experimentes nur drei Mal richtig übersetzt, während die App fünf Mal die inkorrekte Übersetzung *Kérjük* geliefert hat. Da das Wort „Bitte“ nur in Form eines Verbes ins Ungarische übersetzt werden kann – dessen Konjugation verrät, wer bittet und wer gebeten wird (*Kérlek* = ich bitte dich, *Kérem* = ich bitte Sie, *Kérjük* = wir bitten Sie usw.) – kann dieser Fehler auch als Konjugationsfehler verstanden werden, die die nächste Kategorie der Übersetzungsfehler bilden.

8.3.3.2 Konjugationsfehler

Die Konjugation, also die Flexion der Verben nach Person, Numerus, Tempus, Modus, Genus verbi, hat aus Sicht der Kommunikation eine sehr große Bedeutung. Wird ein Verb falsch oder gar nicht flektiert, kann es zu Verständnisproblemen kommen.

Tabelle 31: Transkription des Experimentes, Übersetzungsfehler, Konjugationsfehler (Zeilen 18, 34, 35 und 68)

	Gesprochene Eingabe	ASR + Transkription	MÜ
	A	B	C
18	Ok. Bitte bleiben Sie gleich <u>so stehen</u> , ich werde jetzt an mehreren Stellen <u>eine</u> Ultraschallsonde hinhalten und dabei gleichzeitig auf <u>Ihre</u> Waden drücken <i>Punkt</i> . Es wird Ihnen nichts dabei <u>wehtun</u> .	O. K. Bitte bleiben Sie gleich <u>zustehen</u> , ich werde jetzt an mehreren Stellen <u>in der</u> Ultraschallsonde hinhalten und dabei gleichzeitig auf <u>ihre</u> Waden drücken . Es wird Ihnen nichts dabei <u>wehtuend</u>	Ok. Kérem, maradjon rendben, én most megáll több helyen az ultrahangos szondán, és nyomjuk meg a borjak ugyanabban az időben. Semmi sem fog bántani
34	Danke, mit der Untersuchung sind wir fertig <i>Komma</i> , Sie können sich wieder anziehen und setzen sich zu mir an den Schreibtisch.	Danke mit der Untersuchung sind wir fertig, Sie können sich wieder anziehen und setzen sich zu mir an den Schreibtisch	Köszönöm a vizsgálatot végeztünk, akkor újra felöltözni , és üljön le velem az íróasztalon

	Gesprochene Eingabe	ASR + Transkription	MÜ
	A	B	C
35	Ich werde Ihnen nun anhand einer Grafik ihr Problem an den Beinvenen erklären.	Ich werde Ihnen nun anhand einer Grafik ihr Problem an den Beinvenen erklären	Én most használ egy grafikus megmagyarázni a problémát a láb vénák

68	Nagyon szépen köszönöm a vizsgálatot és a hasznos tanácsot. A gyógyszert be fogom venni, illetve ezt a gyógyszert meg fogom vásárolni és remélem hogy	Nagyon szépen köszönöm a vizsgálatot és a hasznos tanácsot a gyógyszert be fogom venni, illetve ezt a gyógyszert megfogom vásárolni és remélem hogy	Vielen Dank für die Untersuchung und hilfreiche Ratschläge, die ich nehmen und kaufen dieses Medikament und hoffen, dass
----	---	---	--

Während der Untersuchung kam es mehrmals vor, dass in der von der App gelieferten Übersetzung die Verben einfach in der Infinitivform gelassen wurden. In Zeile 34 wurde die Aussage *Sie können sich wieder anziehen* als *akkor újra felöltözni* (= „sich dann wieder anziehen“) übersetzt. Auch in Zeile 68 wurden die Verben in der Übersetzung nicht flektiert: Die Aussage *ich nehmen und kaufen dieses Medikament* (68C) kann zwar trotz der fehlenden Konjugation verstanden werden, hört sich aber äußerst unprofessionell an.

In Zeile 35 wurde das Verb in der Übersetzung nicht wie in der Originalaussage, nach erster Person Singular, sondern nach dritter Person Singular flektiert: *én most használ egy grafikus megmagyarázni* (= „ich jetzt verwendet einen Grafiker um zu erklären“). Wie man feststellen kann, war die falsche Konjugation des Verbes *verwenden* nicht der einzige Übersetzungsfehler, der in dieser Zeile vorkam (*Grafik – Grafiker*). Der Satz wurde also grammatikalisch und lexikalisch inkorrekt ins Ungarische übersetzt, aber trotzdem vom Patienten verstanden.

In Zeile 18 wurde die Futur 1-Form *ich werde ... drücken* in der Imperativform *nyomjuk meg* (= „Drücken wir!“) ins Ungarische übermittelt. Wortwörtlich ist in der Übersetzung *nyomjuk meg a borjak ugyanabban az időben* (= „drücken wir die Kälber gleichzeitig“) zu lesen, was weder einen Sinn ergibt, noch überhaupt annähernd den Inhalt des Originalsatzes widerspiegelt.

8.3.3.3 Englische Ausdrücke

Tabelle 32: Transkription des Experimentes, Übersetzungsfehler, Englische Ausdrücke (Zeilen 10 und 67)

	Gesprochene Eingabe	ASR + Transkription	MÜ
	A	B	C
10	Télen vagy nyáron, ezt nem figyeltem meg, mindegy <i>pont.</i> Ülőmunkát végzek általában, viszonylag sokat ülök <i>pont.</i> Ami aggaszt, ami gondot okoz, hogy a családban az édesapámnál és a nagyszüleimnél is voltak visszértágulati problémák, ezért aggódom egy kicsit <i>pont.</i>	Télen vagy nyáron ezt nem figyeltem meg mindegy. Ülőmunkát végzek, általában viszonylag sokat ülök. Ami aggaszt, ami gondot okoz, hogy a családban az édesapámnál és a nagyszüleimnél is voltak visszértágulati problémák ezért aggódom egy kicsit.	Im Winter oder Sommer habe ich das sowieso nicht gemerkt. Ich mache einen Sit-down-Job , normalerweise siss eitere ich relativ viel. Was mich beunruhigt, was ein Problem ist, ist, dass mein Vater und meine Großeltern Krampfadern in der Familie hatten, also mache ich mir ein wenig Sorgen.
67	Damit die Venen nicht krank werden, achten Sie bitte bei sportlicher Belastung und körperlicher Anstrengung immer darauf, dass sie nicht mit dem Bauch pressen. Dies tut man üblicherweise, um sich die Belastung zu erleichtern. Es macht aber einen Druck auf die Venenklappen , und kann eventuell auch noch später zu Problemen mit den <u>Venen</u> führen.	Damit die Venen nicht krank werden, achten Sie bitte bei sportlicher Belastung und körperliche Anstrengung immer darauf, dass sie nicht mit dem Bauchpressen. Dies tut man üblicherweise, um sich die Belastung zu erleichtern. Es macht aber einen Druck auf die Venenklappen , und kann eventuell auch noch später zu Problemen mit den <u>Venen</u> führen.	Annak megakadályozására, hogy a vénák szerzés beteg, kérjük, mindig győződjön meg arról, hogy nem nyomja a has, ha gyakorlása és a kifejtése. Ez általában történik, hogy enyhítse a törzs. Ez azonban nyomást gyakorol a véna flaps , és problémákat okozhat a vénák még később.

In Zeile 10 wurde das Wort *ülőmunka* (= „Sitzarbeit“) zwar korrekt transkribiert, in der Übersetzung ist der Begriff aber auf Englisch (*Sit-down-Job*) zu finden. Auch in Zeile 67 wurde das Wort *Venenklappen* nur halbwegs ins Ungarische übermittelt. „*Venen*“ wurde zwar richtig als *véna* übersetzt, „*Klappen*“ ist jedoch als *flaps* in der Übersetzung zu finden.

Diese Fehler lassen vermuten, dass die Applikation *SayHi Translate* den AT zuerst ins Englische übersetzt und ihn erst danach in die jeweilige Zielsprache transferiert (falls eine entsprechende Übersetzung vorliegt).

8.3.4 Schwerwiegende Übersetzungsfehler

In diesem Kapitel werden Übersetzungsfehler dargestellt, welche die Übersetzung schwer nachvollziehbar oder sogar vollkommen unverständlich machten.

Die Fehlerursache ist oft schwer auszumachen. Einige schwerwiegende Übersetzungsfehler scheinen meine oben formulierte Annahme, wonach der AT zuerst ins Englische und erst danach in die tatsächliche ZS übersetzt wird, zu bestätigen.

Tabelle 33: Transkription des Experimentes, Schwerwiegende Übersetzungsfehler (Zeilen 18, 44 und 65)

	Gesprochene Eingabe	ASR + Transkription	MÜ
	A	B	C
18	Ok. Bitte bleiben Sie gleich <u>so stehen</u> , ich werde jetzt an mehreren Stellen <u>eine</u> Ultraschallsonde hinhalten und dabei gleichzeitig auf <u>Ihre Waden</u> drücken <u>Punkt</u> . Es wird Ihnen nichts dabei <u>wehtun</u> .	O. K. Bitte bleiben Sie gleich <u>zustehen</u> , ich werde jetzt an mehreren Stellen <u>in der</u> Ultraschallsonde hinhalten und dabei gleichzeitig auf <u>ihre Waden</u> drücken. Es wird Ihnen nichts dabei <u>wehtuend</u>	Ok. Kérem, maradjon rendben, én most megáll több helyen az ultrahangos szondán, és nyomjuk meg a <u>borjak</u> ugyanabban az időben. Semmi sem fog bántani
44	Das ist bei <u>Ihnen</u> aber nicht der Fall <u>Komma</u> , die <u>Klappen</u> funktionieren alle <u>sehr gut</u> .	Das ist bei <u>ihnen</u> aber nicht der Fall, die <u>Klappen</u> funktionieren alle <u>sercan</u>	De nem ez a helyzet velük, a <u>füleket</u> minden munkát sercan
65	Dies belastet die <u>Venenklappen</u> und kann zu Schäden führen, wie sie in der Bezeichnung, die ich jetzt aufgedeckt habe, gezeichnet sind.	Dies belastet die <u>Venenklappen</u> und kann zu Schäden führen wie sie in der Bezeichnung, die ich jetzt aufgedeckt habe gezeichnet sind	Ez teszi a megterhelést a <u>véna füleket</u> és károsíthatja a károkat, mivel azok szerepelnek a kijelölésben már fedezetlen

In Zeile 18 wurde das Wort *Waden* zwar korrekt von der ASR-Komponente transkribiert, ins Ungarische jedoch als *borjak* (= „Kälber“) übersetzt. Grund dafür kann die englische Übersetzung sein: Das deutsche Wort *Waden* wurde vermutlich als „calves“ ins Englische transferiert, „calf“ ist aber mehrdeutig und kann sich sowohl auf „die Wade“ als auch auf „das Kalb“

beziehen. Bei der Übersetzung aus dem Englischen ins Ungarische wurde wahrscheinlich die falsche Bedeutung des Wortes ausgewählt, was das Satzverständnis massiv beeinträchtigt hat.

Ein ähnliches Problem ist auch bei der Übersetzung des Wortes *Venenklappen* zu beobachten. Dass dieser Fachbegriff von der App in allen Fällen richtig transkribiert wurde, ist auf jeden Fall positiv zu bewerten, seine Übersetzung stellte jedoch mehrmals eine Herausforderung für die MÜ-Komponente dar. Sowohl in Zeile 44 als auch in Zeile 65 wurde das Wort *Venenklappen* als *fülek* (= „Ohren“) übersetzt. Die Ursache dieses Fehlers ist erneut in der englischen Übertragung zu finden. Das deutsche Wort *Venenklappen* kann als „vein flaps“ ins Englische übersetzt werden, „flaps“ hat jedoch sehr viele Bedeutungen und kann sich unter anderem auf die Ohren beziehen.

Tabelle 34: Transkription des Experimentes, Schwerwiegende Übersetzungsfehler (Zeile 42)

	Gesprochene Eingabe	ASR + Transkription	MÜ
	A	B	C
42	Damit das Blut aber nur in die Richtung "Herz" und von der <u>Außenvene</u> zur inneren <u>Vene</u> fließen kann, braucht es eine Venenklappen , die das Blut nur in die richtige Richtung fließen lassen.	Damit das Blut aber nur in die Richtung Herz und von der <u>Außenlinie</u> zur inneren <u>Biene</u> fließen kann, braucht es eine Venenklappen die das Blut nur in die richtige Richtung fließen lassen	De a vér áramlását csak az irányba a szív és a külső vonal a belső méh, akkor szüksége van a véna szelep , amely csak lehetővé teszi a vér áramlását a helyes irányba.

Auch die Mehrdeutigkeit des deutschen Wortes „Klappe“ hat sich in der Übersetzung gezeigt: In Zeile 42 wurde *Venenklappen* als *véna szelep* (wortwörtlich zu Deutsch: „*Venen Ventil*“) übersetzt. Diese Übersetzung ergab keinen Sinn und hat die Verständlichkeit des – wegen der Spracherkennungsfehler ohnehin kaum nachvollziehbaren Satzes – weiter erschwert.

Tabelle 35: Transkription des Experimentes, Schwerwiegende Übersetzungsfehler (Zeile 32)

32	Drehen Sie sich jetzt bitte zu mir nach vorne und <u>halten die Hose</u> ein wenig in den Schritt hinein <i>Komma</i> , damit ich bei der Leiste messen kann.	Drehen Sie sich jetzt bitte zu mir nach vorne und <u>halt in die Hose</u> ein wenig in den Schritt hinein, damit ich bei der Leiste messen kann	Most kérem, forduljon hozzám előre, és tartsa a nadrágot egy kicsit az y , hogy tudom mérni a Lágyulás
----	---	---	--

In Zeile 32 gab es zwei schwerwiegende Übersetzungsfehler: Wie bereits in Kapitel 8.2 erwähnt, wurde das Wort *Schritt* einfach mit dem Buchstaben *y* ins Ungarische übersetzt, was den Informationstransfer stark beeinträchtigte. Auch die Übersetzung des Wortes *Leiste* stellte für die App ein Problem dar: Statt „*lágycék*“ (= „*Leiste*“) ist *Lágyulás* (= „*Erweichung*“) in der Übersetzung zu lesen (warum das ungarische Wort großgeschrieben wurde, ist unklar). Die beiden Wörter hören sich zwar im Ungarischen ähnlich an, ihre Bedeutungen unterscheiden sich aber sehr. Der Patient konnte diese Aussage wahrscheinlich nicht gut verstehen, denn er hat – statt seine Hose ein wenig in den Schritt hineinzuhalten – nur sein T-Shirt hochgezogen. In diesem Fall hat es dem Arzt nicht gestört, da er die Messung auch so durchführen konnte.

Tabelle 36: Transkription des Experimentes, Schwerwiegende Übersetzungsfehler (Zeile 58)

	Gesprochene Eingabe	ASR + Transkription	MÜ
	A	B	C
58	Nemrégén kezdtem el ööö... tornaterembe, kondicionáló terembe jágni, remélem az is segíteni fog.	Nemrégén keztem el tudom tornaterembe kondicionáló terembe jágni remélem az is segíteni fog.	Ich habe in der Lage, in ein Fitnessstudio vor langer Zeit zu gehen, so hoffe ich, es hilft.

In Zeile 58 wurde das Wort *nemrégén* (= „*kürzlich*“, „*vor kurzem*“) falsch übertragen: In der deutschen Übersetzung ist genau das Gegenteil – *in ein Fitnessstudio vor langer Zeit zu gehen* – zu lesen. Inwiefern die Spracherkennungsfehler anderer Wörter in diesem Satz für diese Fehlübersetzung verantwortlich waren, kann schwer bestimmt werden. Ob der Patient seit langem das Fitnessstudio besucht oder erst vor kurzem mit dem Trainieren begonnen hat, spielt in dieser speziellen Situation keine wesentliche Rolle. Würde sich die Frage jedoch z.B. im Rahmen eines Aufklärungsgespräches vor der Operation darauf beziehen, seit wann der Patient ein bestimmtes Medikament nimmt, wäre es möglicherweise von existenzieller Bedeutung, ob er damit vor *langer* oder *kurzer* Zeit angefangen hat.

Tabelle 37: Transkription des Experimentes, Schwerwiegende Übersetzungsfehler (Zeilen 15, 36, 51 und 64)

	Gesprochene Eingabe	ASR + Transkription	MÜ
	A	B	C
15	Bevor ich Ihnen das <u>beantworte</u>, möchte ich Sie gerne untersuchen <i>Punkt</i>. Dazu werde ich mir ihre Beine im Stehen anschauen und einen Ultraschall machen <i>Punkt</i>. Dazu müssen Sie bitte in der Garderobe hinter Ihnen die Beine <u>freimachen</u> <i>Punkt</i>. Bitte ziehen Sie Ihre Hose, die Schuhe und die Strümpfe aus und stellen Sie sich auf das kleine <u>Podest</u> <i>Komma</i>, wo die Füße grün eingezeichnet sind.	Bevor ich Ihnen das <u>beantwortet</u> möchte ich Sie gerne untersuchen. Dazu werde ich mir ihre Beine im Stehen anschauen und einen Ultraschall machen. Dazu müssen Sie bitte in der Garderobe hinter ihnen die Beine <u>freimachen</u>. Bitte ziehen Sie Ihre Hose die Schuhe und die Strümpfe aus und stellen Sie sich auf das kleine <u>Protest</u>, wo die Füße grün eingezeichnet sind	Mielőtt válaszolnék nektek, hogy szeretném megvizsgálni téged. Fogom nézni a lábát állva, és nem egy ultrahang. Ehhez kérjük, hogy <u>törölje</u> a lábát az öltözőben mögöttük. Kérjük, vegye le a nadrágját, és harisnya és állni a kis tiltakozás, ahol a lábad készült zöld
36	Diese Zeichnung ist eine vereinfachte Darstellung der Venen des Beines <i>Punkt</i>. Man muss sich die blauen Linien räumlich als Röhren <u>vorstellen</u> <i>Punkt</i>.	Diese Zeichnung ist eine vereinfachte Darstellung der Venen des Beines. Man muss sich die blauen Linien räumlich als Röhren <u>vorstellen</u>.	Ez a rajz egy egyszerűsített ábrázolása a vénák a lábát. <u>Meg kell gondolni</u>, a kék vonalak a csöveket.
51	Dieses Problem können Sie nur lösen, indem sie während <u>Ihrer</u> Berufstätigkeit immer wieder <u>versuchen</u> <i>Komma</i>, <u>aufzustehen</u> und kurze Wege zu Fuß zu erledigen <i>Punkt</i>. Weiters sollten Sie in der Freizeit sich bemühen, viel Bewegung mit den Beinen zu machen <i>Punkt</i>. Auch kann ich Ihnen ein Medikament verschreiben, dass dieses Problem mildern wird.	Dieses Problem können Sie nur lösen, indem sie während <u>ihrer</u> Berufstätigkeit immer wieder <u>versuchen</u>, <u>aufzustehen</u> und kurze Wege zu Fuß zu erledigen. Weiters sollten Sie in der Freizeit sich bemühen, viel Bewegung mit den Beinen zu machen. Auch kann ich Ihnen ein Medikament verschreiben, dass dieses Problem mildern wird	Csak akkor tudja megoldani ezt a problémát, ha többször <u>próbál feljutni</u>, és gyalog rövid távolságok munka közben. Azonkívül, Önnek kellene csinál egy erőfeszítés-ban-a szabad idő-hoz csinál sok gyakorlás-val-a lábak. Is tudok felírni neked egy gyógyszer, amely enyhíteni ezt a problémát

	Gesprochene Eingabe	ASR + Transkription	MÜ
	A	B	C
64	Beim Husten oder <u>Pressen</u> mit dem Bauch wird mit großem Druck das Venenblut wieder in die Beine <u>zurückgedrückt</u> .	Beim Husten oder <u>pressend</u> mit dem Bauch wird mit großem Druck das Venenblut wieder in die Beine <u>zurückgedrückt</u>	Amikor köhöt vagy présel a hasüregbe, a véna vére nagy nyomással <u>visszateredik</u> a lábakra

Auch bei der Übersetzung einiger Verben können schwerwiegende Fehler beobachtet werden.

In Zeile 15 wurde *die Beine freimachen* als *töröje a lábát* (= „wischen Sie Ihre Beine“) übersetzt, was zu Einschränkungen beim Verständnis führte.

In Zeile 36 ist das Verb *vorstellen* als *meg kell gondolni* (= „man muss es bedenken“) in der Übersetzung zu lesen, was zwar grundsätzlich nicht zur massiven Sinnverschiebung geführt hat, trotzdem sehr störend ist.

Der Ratschlag, wonach der Patient versuchen sollte, während seiner Berufstätigkeit immer wieder aufzustehen (51A), wurde ebenfalls falsch in die ZS übertragen. Es wurde ihm in der Übersetzung (51C) empfohlen, „immer wieder versuchen, heraufzukommen“ (*többször próbál feljutni*); was wenig Sinn ergab.

In Zeile 61 wurde das Verb *zurückgedrückt* statt mit „*visszanyomódik*“ als *visszateredik* übersetzt. Diese Übersetzung ist auch deshalb interessant, weil dieses Verb im Ungarischen gar nicht existiert. Vielleicht wollte das Programm das Wort „*visszaterelődik*“ (= „zurückgetrieben“) schreiben, was zwar nicht vollständig dem Originalausdruck entsprochen hätte, trotzdem besser verständlich gewesen wäre.

8.4 Usability

Wie bereits in Kapitel 7.4.2 erwähnt, wird unter Usability in dieser Arbeit die „Natürlichkeit der Interaktion“ verstanden. Grundsätzlich geht es bei diesem Aspekt darum, wie sich die Kommunikation zwischen zwei Menschen dadurch verändert, dass die verbale Verständigung lediglich mit Hilfe eines Tablets und der entsprechenden Software ermöglicht wird.

In der Arbeit wurde mehrmals die Wichtigkeit der Schaffung einer möglichst realitätsnahen Situation betont. Aus diesem Grund hat der Patient dem Arzt die App persönlich vorgestellt und sich mit wirklichen Beschwerden an ihn gewendet. Auf ein Problem, das schon vor dem Experiment auftauchte, soll nochmals ausdrücklich hingewiesen werden. Es geht um die Frage der Internetverbindung. Da der Patient auf seinem Tablet keinen Internetzugang mittels SIM-Card hatte – die App jedoch offline nicht verwendbar ist – musste nach dem WLAN-Passwort gefragt werden.¹⁴ Da der Patient wegen seiner mangelnden Sprachkenntnisse ohne die App nicht hätte erklären können, dass und wozu er das Passwort benötigt, musste ich an seiner Stelle beim Arzt nachfragen.

Ein anderer wichtiger Punkt, der behandelt werden muss, ist der Zeitfaktor. Wie bereits erwähnt, dauerte das Experiment 28 Minuten. Die Länge der Untersuchung wurde auch vom Arzt im Interview hervorgehoben:

Es könnte natürlich schneller gehen, üblicherweise würde ich für den Patienten eine Viertelstunde brauchen, so brauche ich doppelt so lang. (IA: 4)

Die Tatsache, dass die Untersuchung doppelt so lange dauerte, wurde jedoch nicht durch die Verwendung des Tablets, sondern durch das Dolmetschen des Gesprächs verursacht. Auch mit HumandolmetscherInnen hätte der Arzt länger für die Untersuchung gebraucht, da beim Konsekutivdolmetschen das Gesagte immer wiederholt wird. Man darf aber nicht vergessen, dass die Verwendungsweise der App zuerst vorgestellt werden musste und dass die Transkription während der Untersuchung mehrmals manuell korrigiert wurde, was natürlich Zeit kostete. Auch das Hin- und Herdrehen des Tablets hat den Gesprächsfluss beeinträchtigt. Außerdem war die Tatsache, dass *SayHi Translate* im Konsekutivmodus „dolmetschte“, in erster Linie bei

¹⁴Natürlich hatte ich vor dem Experiment nachgefragt, ob ein stabiles kabelloses Netzwerk in der Ordination vorhanden ist, da dies aus Sicht meines Experimentes sehr wichtig war. Im wirklichen Leben können PatientInnen aber nicht immer davon ausgehen, dass ein Router in der Ordination eingerichtet ist und sie auch das Passwort bekommen.

der Untersuchung mit der Ultraschallsonde sehr störend, was auch durch die nonverbale Kommunikation des Arztes ersichtlich wurde:

Tabelle 38: Transkription des Experimentes, nonverbale Kommunikation des Arztes während der Untersuchung

	Gesprochene Eingabe	ASR + Transkription	MÜ	Bildbeschreibung, nonverbale Kommunikation
	A	B	C	D
18	Ok. Bitte bleiben Sie gleich <u>so stehen</u> , ich werde jetzt an mehreren Stellen <u>eine</u> Ultraschallsonde hinhalten und dabei gleichzeitig auf <u>Ihre</u> Waden drücken <i>Punkt</i> . Es wird Ihnen nichts dabei <u>wehtun</u> .	O. K. Bitte bleiben Sie gleich <u>zustehen</u> , ich werde jetzt an mehreren Stellen <u>in der</u> Ultraschallsonde hinhalten und dabei gleichzeitig auf <u>ihre</u> Waden drücken. Es wird Ihnen nichts dabei <u>wehtuend</u>	Ok. Kérem, maradjon rendben, én most megáll több helyen az ultrahangos szondán, és nyomjuk meg a borjak ugyanabban az időben. Semmi sem fog bántani	<i>der Arzt möchte mit der Untersuchung anfangen, muss aber auf die Übersetzung warten, da aber die App hier ein wenig länger braucht, zeigt er ungeduldig auf das Tablet</i>
27	Bitte <u>stellen</u> Sie jetzt das linke Bein mit einem halben Schritt Weg nach vorne weg <i>Punkt</i> .	<u>Wählen</u> Sie jetzt das linke Bein mit einem halben Schritt Weg nach vorne weg.	Most válassza ki a bal lábát egy fél lépéssel el.	<i>Während der Arzt auf die Übersetzung wartet, macht er kreisende, ungeduldige Bewegungen mit den Händen. Er schiebt das linke Bein des Patienten einen halben Schritt nach vorne.</i>



Abbildung 21: Nonverbale Kommunikation des Arztes während der Untersuchung (Ausschnitt aus dem Video)

Die Ungeduld des Arztes lässt darauf schließen, dass er sich wahrscheinlich eine simultane Übertragung gewünscht hätte. HumandolmetscherInnen würden in einer solchen Situation vermutlich das Flüsterdolmetschen wählen, um die Untersuchung zu beschleunigen.

Die nonverbale Kommunikation spielte während der ganzen Untersuchung eine sehr wichtige Rolle.

Tabelle 39: Transkription des Experimentes, nonverbale Kommunikation (Zeilen 1 – 4, 52 – 53)

	Gesprochene Eingabe	ASR + Transkription	MÜ	Bildbeschreibung, nonverbale Kommunikation
	A	B	C	D
1	Hallo!	x	x	<i>gibt die Hand</i>
2	Bitte, setzen Sie sich!	x	x	<i>schüttelt die Hand und zeigt auf den Stuhl</i>

3	Jó napot kívánok doktor úr, köszönöm szépen, hogy időt szánt rám <i>pont</i> . Bécsben dolgozom, de sajnos nem beszélek németül, ezért hoztam el ezt az applikációt, ami remélem segít nekünk a kommunikációban <i>pont</i> .	Napot kívánok doktor úr köszönöm szépen hogy időt szánt rám. bécsben dolgozom, de sajnos nem beszélek németül ezért hoztam el ezt az applikációt amit remélem segít nekünk majd a kommunikációban.	Ich wünsche Ihnen einen Tag, Doktor. Ich arbeite in Wien, aber leider spreche ich kein Deutsch, deshalb habe ich diese App mitgebracht, die uns hoffentlich helfen wird, zu kommunizieren.	<i>setzt sich an den Schreibtisch und beginnt zu sprechen; der Arzt möchte gleich antworten und greift nach dem Tablet, der Patient erhebt jedoch den Zeigefinger, um zu signalisieren, dass er noch etwas sagen möchte</i>
4	Gyorsan elmagyarázom, hogy hogyan működik a program. Itt vannak a nyelvek, amikor beszélünk, ezeket a gombokat kell megnyomni és utána kell beszélni <i>pont</i> .	Gyorsan elmagyarázom hogy hogyan működik a program, itt vannak a nyelvek, amikor beszélünk ezeket a gombokat kell megnyomni és utána kell beszélni.	Ich erkläre schnell, wie das Programm funktioniert, hier sind die Sprachen, wenn wir über diese Tasten sprechen, die gedrückt werden und dann sprechen.	<i>zeigt während der Spracheingabe auf die Tasten; nachdem das Gerät übersetzt hat, dreht der Arzt das Tablet zu sich</i>

52	Haben Sie zu dieser Erklärung irgendwelche Fragen?	Haben Sie zu dieser Erklärung irgendwelche Fragen	Kérdése van ezzel a nyilatkozattal kapcsolatban?	
53				<i>schüttelt den Kopf</i>

Wie man in den Zeilen 1, 2 und 53 feststellen kann, musste während der Untersuchung nicht immer das Tablet verwendet werden, um kommunizieren zu können: Händeschütteln, Nicken, Kopfschütteln usw. sind Signale, die – zumindest in unserem Kulturkreis – auch ohne Dolmetschen verstanden werden können. Auch die Grafik, anhand deren die Krankheit der Beinvenen vom Arzt erklärt wurde, war eine große Hilfe für den Patienten. Obwohl er der Erklärung nicht immer folgen konnte, konnte er sich gut vorstellen, worüber der Arzt spricht.



Abbildung 22: Körpersprache während des Experimentes; Händeschütteln bei der Begrüßung; Ende der Untersuchung: Patient nickt, Arzt hält Daumen nach oben (Ausschnitt aus dem Video)



Abbildung 23: Der Arzt verwendet eine Grafik, um die Krankheit an den Beinvenen zu erklären (Ausschnitt aus dem Video)

Wie bereits in Kapitel 8.2 dargestellt, wurde das Wort „husten“ von der ASR-Komponente immer falsch erkannt und deswegen auch fehlerhaft übersetzt. Dadurch entstand während der Untersuchung ein Kommunikationsproblem, da der Patient nicht verstand, dass er husten sollte. Der Arzt erkannte jedoch das Problem und hustete selbst, um zu zeigen, was gemeint war. Der Patient verstand dies und hustete nach dem Vorbild des Arztes. Danach signalisierte der Arzt – statt die Applikation zu verwenden – immer mit einem Husten seinerseits, dass er dasselbe vom Patienten erwartete. Das Problem konnte also nonverbal, ohne die App gelöst werden.

Auch der Blickkontakt der Teilnehmer kann als ein wichtiges Ausdrucksmittel der Körpersprache verstanden werden. Bei diesem Experiment war gut zu beobachten, dass die Gesprächspartner – statt sich gegenseitig anzuschauen und den Blickkontakt während des Sprechens zu halten – meistens den Bildschirm des Tablets beobachteten. Dies kann natürlich zum Teil dadurch erklärt werden, dass die Teilnehmer keine gemeinsame Sprache besaßen und ausschließlich auf die Leistung der Applikation angewiesen waren. Man kann auch bei humanen Dolmetscheinsätzen den Eindruck bekommen, dass die GesprächspartnerInnen eher mit den DolmetscherInnen kommunizieren, da sie während des Sprechens natürlich verstanden werden möchten.

Professionelle DolmetscherInnen sind in der Lage, die Wörter der Kommunikationsparteien so wiederzugeben, dass auch Stil und Emotionalität der Aussagen erhalten bleibt. Das schaffte diese Applikation eindeutig nicht: Die Sprachausgabe war monoton und sie klang mechanisch. Abgesehen von den vorher erwähnten Stellen im Gespräch, in denen die Körpersprache die Verständigung stark unterstützte (Nicken, Kopfschütteln usw.), kann man feststellen, dass die Nutzung des Tablets die „Natürlichkeit“ der menschlichen Kommunikation negativ beeinflusst hat. Sowohl der Arzt als auch der Patient sprachen langsam und artikuliert, was eine negative Auswirkung auf die Flüssigkeit der Kommunikation hatte, die Aussagen waren frei von Emotionen und wirkten ein wenig abgebrochen.

Ein weiterer interessanter Punkt ist die Positionierung des Tablets während des Experimentes. Als der Patient in den Ordinationsraum eintrat, hatte er es in der Hand. Nachdem er Platz genommen hatte, legte er das Gerät so auf den Schreibtisch, dass auch der Arzt es sehen konnte. Nachdem er sich selbst und die Applikation vorgestellt hatte, überreichte er das Tablet dem Arzt, sodass dieser antworten konnte. Während des Arzt-Patienten-Gespräches drehten die Teilnehmer das Gerät auf dem Schreibtisch hin und her, was ein wenig mühsam schien. Als der Arzt die Untersuchung mit der Ultraschallsonde durchführte, legte er das Tablet auf seinen Schoß. Da er während und nach der Untersuchung viel mehr sprach als der Patient, war er derjenige, der das Tablet mit sich nahm, als sie sich im Raum bewegten. Die gute und flexible

Handhabung des Tablets während des Experiments ist wahrscheinlich der Tatsache zu verdanken, dass sowohl der Arzt als auch der Patient viel Erfahrung im Umgang mit mobilen Endgeräten hatten.

Da die Teilnehmer des Experimentes selbst entscheiden konnten, wie lange sie sprechen, wurden die zu übersetzenden Textsegmente oft sehr lang. Aus diesem Grund brauchte auch die Applikation länger, um die Texte zu verarbeiten und zu übersetzen, was zu Pausen in der Kommunikation führte. Das Warten auf die Übertragung wurde vom Patienten auch im Rahmen des Interviews negativ erwähnt:

Die Textverarbeitung der App führte tatsächlich oft zum „Leerlauf“ in der Kommunikation. Meiner Ansicht nach war dies ein bisschen störend, da das Gespräch so immer abgebrochen wurde, ich finde die App aber grundsätzlich verwendbar. (IP: 3)

Der Arzt empfand diese Pausen nicht als sehr störend, im Interview bemerkte er nur, dass es natürlich schneller gehen könnte (vgl. IA: 4).

Die falsche Transkription des Gesagten war ein Problem, das in der Arbeit schon mehrmals hervorgehoben wurde. Wie man in Kapitel 8.2 lesen kann, kam es trotz der großen Anzahl der Spracherkennungsfehler nur dreimal vor, dass die Transkription manuell korrigiert wurde.

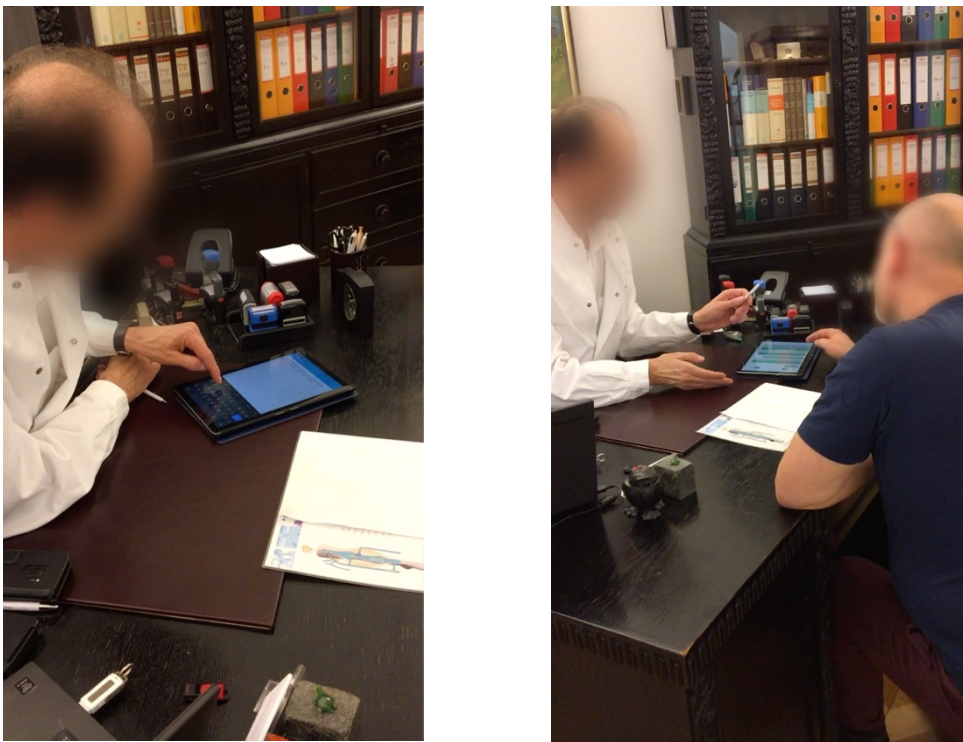


Abbildung 24: Manuelle Korrekturfunktion; der Patient zeigt dem Arzt, wie man die Transkription bearbeiten kann; der Arzt korrigiert den Text mit Hilfe einer Tastatur (Ausschnitt aus dem Video)

Der Arzt fand die Korrektur grundsätzlich nicht störend, vielmehr aber, dass die App ihn oft falsch verstand (vgl. IA: 3). Der Patient äußerte sich diesbezüglich folgenderweise:

Die manuelle Korrektur war gar nicht störend, sie machte aber die ärztliche Untersuchung irgendwie wirklichkeitsfremd. (IP: 6)

Die ungarische Übersetzung war oft sowohl grammatikalisch als auch lexikalisch fehlerhaft, es kam jedoch während des Gesprächs nur einmal vor, dass der Patient nachfragte. Es stellt sich die berechnigte Frage, warum er nicht öfter zum Ausdruck brachte, dass er das Gesagte nicht verstanden hatte.

Ja, wie gesagt, gab es Stellen, wo ich den Arzt tatsächlich nicht verstanden habe. Da ich aber schon ein wenig Erfahrung mit der App hatte, wusste ich, dass sie das Gesagte nicht immer gut versteht und manchmal Fehlübersetzungen liefert. Der Arzt hat sehr langsam gesprochen, ich wusste, dass nicht er für die schlechten Übersetzungen verantwortlich ist, sondern die Applikation selbst. Er hätte keinen Sinn gehabt, nachzufragen. Vor allem, weil wir nicht z.B. einen wichtigen Operationstermin vereinbart haben, er hat „lediglich“ Tipps für eine gesunde Lebensführung gegeben. (IP: 5)

Da der Patient nicht zurückfragte, dachte der Arzt, dass seine verbalen Informationen korrekt ins Ungarische transferiert wurden und dass der Patient alles gut verstehen konnte:

- *Hatten Sie den Eindruck, dass der Patient Sie gut verstanden hat, oder hat er nur so getan, als ob er alles gut verstanden hätte?*
- *Also ich glaube schon, dass er alles gut verstanden hat. (IA: 5)*

Es gab einen weiteren Störfaktor, der erwähnt werden muss: das sind die Mitteilungen diverser anderer Apps, die während des Experimentes mittels Signalton auf dem Tablet ankamen. Da eine stabile Internetverbindung vorhanden war und diese Mitteilungssignale nicht ausgeschaltet worden waren, erhielt der Patient während der Untersuchung Nachrichten von mehreren Social Media Plattformen, deren Vorschau auf dem Bildschirm eingeblendet wurde. Solche unerwarteten Mitteilungen können zum einen den Sprecher ablenken und zum anderen kann der Nachrichtenton die automatische Spracherkennung stören. Im Laufe des Experimentes empfing der Patient nur drei Mitteilungen, die den Kommunikationsablauf glücklicherweise nicht beeinträchtigten. Im Idealfall sollten Benachrichtigungen anderer Apps also bei der Verwendung der Dolmetsch-App ausgeschaltet werden.

Trotz der Störfaktoren und der großen Anzahl der Fehlinterpretationen durch die App lässt sich feststellen, dass der Zweck der Untersuchung im Großen und Ganzen erfüllt werden konnte, was der Patient im Interview auch betonte:

Ich fand den Arzt sehr sympathisch, ich hatte den Eindruck, dass er trotz der schwierigen Kommunikation mein Grundproblem verstanden hat. Er führte die Untersuchung durch und konnte mir eine beruhigende Antwort bezüglich einer eventuell erblichen Krankheit geben. Er war fachlich sehr authentisch, hätte ich erstere Probleme, würde ich mich gerne an ihn wenden. Nach der Untersuchung habe ich sogar eine anatomische Aufklärung bzw. Tipps für eine gesunde Lebensführung bekommen, die ich aber in den meisten Fällen wegen der Fehlübersetzungen nicht wirklich verstanden habe. (IP: 7)

Dass das Experiment aufgrund von technischen Hindernissen oder Verständigungsproblemen nicht unterbrochen oder völlig abgebrochen werden musste, ist auf jeden Fall positiv anzumerken.

9. Zusammenfassung und Schlussfolgerungen

Ziel des Experimentes war einerseits, eine Antwort auf die Frage zu geben, inwieweit die mobile Dolmetsch-Applikation *SayHi Translate* dafür geeignet ist, die Kommunikation im Rahmen eines Arzt-Patient-Gesprächs im Sprachenpaar Deutsch-Ungarisch zu ermöglichen; andererseits herauszufinden, inwiefern die Teilnehmer des Experimentes selbst mit der Leistung der App zufrieden waren und ob sie auch in Zukunft in einer ähnlichen Situation auf die Applikation zurückgreifen würden.

Anhand des von mir durchgeführten Experimentes und dessen Evaluierung können folgende Schlussfolgerungen gezogen werden:

Die Applikation *SayHi Translate* ist grundsätzlich transparent aufgebaut und einfach zu bedienen. Sie verfügt über eine Vielzahl von Sprachen und Dialekten, die bezüglich der Stimme und Geschwindigkeit der Sprachausgabe individualisiert werden können. Es spricht außerdem für die App, dass sie kostenlos auf alle iOS- und Android-Geräte heruntergeladen werden kann. Trotz der Einschränkung, dass zu ihrer Verwendung eine stabile Internetverbindung notwendig ist, ist *SayHi Translate* aus Sicht der BenutzerInnenfreundlichkeit auf jeden Fall positiv zu bewerten. Doch obwohl eine gute Benutzeroberfläche die Handhabung erleichtert, sie allein ist nur ein kleiner Bestandteil der App. Für eine Gesamtbeurteilung muss das Zusammenspiel der verschiedenen Komponenten der Applikation betrachtet werden.

Um die Leistung von *SayHi Translate* zu bewerten, wurden zwei Fehleranalysen durchgeführt: einerseits wurden die Spracherkennungsfehler untersucht, andererseits wurden die Übersetzungsfehler, die trotz der richtigen Transkription vorkamen, gewichtet und analysiert.

Durch die Evaluierung der Spracherkennungsfehler wurde sichtbar, dass die falsche Transkription des Gesagten oft für unverständliche Übersetzungen verantwortlich ist. Insgesamt gab es während des Experimentes 53 Fehler in der automatischen Spracherkennung, 29 davon beeinträchtigten sogar die lexikalische Korrektheit bzw. den Informationstransfer der Übersetzung stark (z.B. *husten* – *Houston*, *Podest* – *Protest*, *Sie* – *sie*). Ich möchte jedoch nicht unerwähnt lassen, dass diese 53 Fehler im Prozentsatz ausgedrückt ca. 4,5% der gesprochenen und transkribierten Wörter ausmachen (insgesamt gab es im Gespräch 1171 Wörter). Würde man nur die 95,5% der korrekt transkribierten Wörter in Betracht ziehen, so wäre die App wahrscheinlich auch aus Sicht der automatischen Spracherkennung positiv zu bewerten. Allerdings darf auf keinen Fall übersehen werden, dass schon ein einziges falsch erkanntes Wort einen ganzen Satz unverständlich machen kann. Da ich bei der Evaluierung des Experimentes aufgabenbezogen vorgehe, ist in erster Linie nicht die Anzahl der Fehler ausschlaggebend,

sondert ihre Wirkung auf den Ablauf und auf die Qualität der Kommunikation. Die Spracherkennungsfehler spiegelten sich stark in den Übersetzungen wider, der Patient verlor während des Aufklärungsgesprächs mehrmals den Faden, da die Applikation „nicht einmal annähernd das übersetzt hat“ (IP: 4), worüber der Arzt sprach.

Auch die große Anzahl der Übersetzungsfehler übte eine negative Wirkung auf den Informationstransfer zwischen den zwei Sprachen aus. Man sieht, dass die Übersetzungen aus dem Ungarischen ins Deutsche viel verständlicher sind als umgekehrt. In den deutschen Übersetzungen sind zwar kleinere stilistische, syntaktische oder grammatikalische Fehler zu finden, die übersetzten Textsegmente sind jedoch grundsätzlich gut zu verstehen. Bei den ungarischen Übersetzungen kann jedoch eine große Anzahl von Übersetzungsfehlern beobachtet werden. An dieser Stelle möchte ich noch einmal betonen, dass die falsche Wortstellung und Präpositionen, die fehlerhafte Grammatik bzw. die Deklinationsfehler vermutlich oft durch die falsche automatische Segmentierung der Texteinheiten entstanden sind.

Um entscheiden zu können, inwieweit *SayHi Translate* eine gelungene Arzt-Patient-Kommunikation ermöglichen konnte, muss außerdem untersucht werden, ob das Kommunikationsziel beider Gesprächspartner erfüllt werden konnte.

Bezüglich der Anwendbarkeit der App haben Arzt und Patient sehr unterschiedliche Meinungen formuliert. Der Arzt beantwortete die Frage nach einer möglichen weiteren Nutzung der App in seinem beruflichen Kontext folgenderweise:

- *Falls jemand in der Zukunft mit dieser App kommen würde, wie würden Sie reagieren? Würden Sie auf einen Dolmetscher bestehen oder es mit der App probieren?*
- *Ich würde mir mehr Zeit einteilen und würde es probieren mit der App.* (IA: 6)

Dass der Arzt auch in der Zukunft auf die App zurückgreifen würde, ist grundsätzlich damit zu erklären, dass sein Kommunikationsziel erfüllt wurde: Da die Antworten des Patienten relativ gut und verständlich ins Deutsche „gedolmetscht“ wurden, bekam er alle relevanten Informationen über die Beschwerden. Er konnte die Untersuchung mit der Ultraschallsonde durchführen, anhand der Ergebnisse eine Diagnose erstellen und den Patienten beruhigen, dass er die Krankheit glücklicherweise nicht geerbt hat, sowie ein Medikament verschreiben, das ihm in der Prävention helfen kann. Da sich der Patient für die hilfreichen Ratschläge bedankt und keine Fragen gestellt hatte, vermutete der Arzt, dass er tatsächlich alles verstanden hat. Der Patient gab aber im Interview an, dass er wegen der Fehlübersetzungen mehrmals den Faden verloren hat.

[...] das Warten auf die Übersetzung war für mich nicht störend, wegen der Fehlübersetzungen habe ich aber leider mehrmals den Faden verloren, da das Gerät nicht einmal annähernd das übersetzt hat, worüber wir gesprochen haben. Hätte ich ernstere Beschwerden gehabt, hätte ich wahrscheinlich nicht verstanden, was der Arzt mir mitteilen wollte. (IP: 4)

[...] Nach der Untersuchung habe ich sogar eine anatomische Aufklärung bzw. Tipps für eine gesunde Lebensführung bekommen, die ich aber in den meisten Fällen wegen der Fehlübersetzungen nicht wirklich verstanden habe. (IP: 7)

Theoretisch wurde auch das Kommunikationsziel des Patienten erfüllt, da er einerseits ärztliches Verständnis und Gehör für seine Beschwerden, andererseits beruhigende Auskünfte bekam, die er auch verstehen konnte. Trotzdem würde der Patient in einer ähnlichen Situation in Zukunft die App nicht in Anspruch nehmen:

[...] Ich würde das Tablet bzw. diese App nicht gerne in einer solchen Situation in Anspruch nehmen. Natürlich ist es völlig anders, wenn man ein geselliges Gespräch führt oder im Ausland nach Informationen fragen möchte, aber im Fall von Gesundheitsproblemen, wo die genaue und präzise Beschreibung der Symptome für eine Diagnose unerlässlich ist, hätte ich mir viel lieber einen Dolmetscher oder eine Dolmetscherin gewünscht. (IP: 6)

Die negative Einschätzung der App seitens des Patienten lässt sich in erster Linie mit der niedrigen Übersetzungsleistung von *SayHi Translate* erklären. Es darf nicht unerwähnt bleiben, dass die meiste Zeit das Wort vom Arzt ergriffen (81% = 954 Wörter) wurde und nur ca. 19% des Gespräches (215 Wörter) in der Richtung HU → DE erfolgte. Insgesamt gab es 54 übersetzte Texteinheiten, davon wurden 42 ins Ungarische und nur 12 in die deutsche Sprache übertragen. Ausschlaggebend ist außerdem die Tatsache, dass während ca. 50% der deutschen Übersetzungen grammatikalisch korrekt und leicht zu verstehen waren (6 aus 12 Texteinheiten), dieser Wert bei den ungarischen Übersetzungen nur ca. 7% beträgt (3 aus 42 Texteinheiten). Unter den deutschen Zieltexten gab es keinen einzigen, der sehr schwer oder gar nicht verständlich gewesen wäre, bei den ungarischen Übersetzungen betrug der Anteil solcher Texteinheiten ca. 45% (19 aus 42). Dass das Kommunikationsziel der Teilnehmer während der ärztlichen Untersuchung trotz der hohen Fehlerquote der App erfüllt werden konnte, ist vor allem zwei Faktoren zu verdanken: Einerseits der Tatsache, dass die meisten unverständlichen Übersetzungen nicht während der Besprechung der Beschwerden vorkamen, sondern nach der Ultraschalluntersuchung, als der Arzt „lediglich“ die Krankheit erklärte und Tipps für eine gesunde Lebensführung gab. Andererseits spielte auch die nonverbale Kommunikation der Teilnehmer eine

wesentliche Rolle. Die verstärkte Körpersprache trug sehr oft dazu bei, dass sich die Kommunikationsparteien gegenseitig verstehen konnten.

Zusammenfassend lässt sich sagen, dass die Applikation *SayHi Translate* im Sprachenpaar Deutsch-Ungarisch – sowohl bei der automatischen Spracherkennung als auch bei der maschinellen Übersetzung – eine sehr hohe Fehlerquote aufweist. Die Kommunikation wurde zwar grundsätzlich ermöglicht, wichtige Informationen gingen jedoch verloren. Man darf nicht vergessen, dass Fehlinterpretationen im Gesundheitswesen schwerwiegende rechtliche und gesundheitliche Konsequenzen haben können. Für eine genaue Anamnese, Diagnose und Behandlung genügen die klinischen Untersuchungen an sich nicht; eine fachlich korrekte und vollständige Wiedergabe der verbalen Informationen ist ebenfalls essenziell. Ich kann mir gut vorstellen, dass *SayHi Translate* in alltäglichen Situationen (z.B. Small Talk, Reisen) eine große Hilfe darstellen kann, vor allem, wenn einfachere und kürzere Gesprächssequenzen ausgesprochen werden. Im medizinischen Bereich jedoch – und hier möchte ich mich auf den Gedanken von Mitrovic (2017) stützen – kann die Anwendung von MD-Apps vor allem bei ernsten Erkrankungen oder bei lebensrettenden Eingriffen sogar als äußerst gefährlich eingestuft werden.

Das in dieser Arbeit durchgeführte Experiment und die dargestellten Ergebnisse rechtfertigen die Aussage, dass *SayHi Translate* zum jetzigen Zeitpunkt nicht für die einwandfreie Kommunikation im medizinischen Bereich geeignet ist und professionelle DolmetscherInnen im Gesundheitswesen zurzeit nicht ersetzen kann. Es wäre auf jeden Fall wünschenswert, dass sowohl an der Optimierung von Dolmetsch-Apps als auch an finanziellen und rechtlichen Regelungen hinsichtlich Humandolmetschen im Sozial- und Gesundheitsbereich gearbeitet wird. Abschließend möchte ich anmerken, dass die Durchführung und Auswertung ähnlicher Experimente in anderen Sprachenpaaren eine durchaus lohnenswerte Aufgabe für zukünftige Untersuchungen wäre, um ein umfassendes Bild über die Leistung von Dolmetsch-Applikationen in medizinischen Settings zu bekommen.

Bibliographie

- ALPAC (Organization) (1966). *Languages and machines: computers in translation and linguistics: a report*. Washington: National Academy of Sciences, National Research Council.
- Angelelli, Claudia V. (2004). *Medical Interpreting and Cross-cultural Communication*. Cambridge: Cambridge University Press.
- ATS Sprachendienst (2004). Ungarische Grammatik. Die Fälle. <https://www.ats-group.net/deutsch/ungarisch-grammatik.html> (Stand: 13.07.2020)
- Bennet, S. & Gerber, L. (2003): Inside commercial machine translation. In: Somers, H.L. (Hg.) *Computers and translation: a translator's Guide*. Amsterdam: John Benjamins Publishing, 176-190.
- Box for healthcare. <https://www.box.com/de-de/industries/healthcare/hipaa-compliance> (Stand: 14.07.2020)
- Breitenbach, Sina (2018). Meilensteine der Machine Translation: Teil 2 - Schachbretter und neuronale Netze. <https://www.lengoo.de/blog/meilensteine-der-machine-translation-teil-2-schachbretter-und-neuronale-netze/> (Stand: 14.07.2020)
- Brizar, Ivona (2014). *Speech-to-Speech Translation Apps: ein kritischer Vergleich*. Masterarbeit, Universität Wien.
- Bub, Udo (1999). *Anwendungsspezifische Online-Anpassung von Hidden-Markov-Modellen in automatischen Spracherkennungssystemen*. München: Herbert Utz Verlag.
- Cap, Müslüm (2003). *Maschinelle Übersetzung auf dem Prüfstand. Die Evaluierung von Personal Translator 2002 Office Plus Englisch*. Lübeck: Verlag Schmidt-Römhild.
- Eberle, K. (2008). Integration von regel- und statistikbasierten Methoden in der maschinellen Übersetzung. *JLCL – Journal for Language Technology and Computational Linguistics* 24 (3), 37-70.

- Fünfer, Sarah (2013). *Mensch oder Maschine?: Dolmetscher und maschinelles Dolmetschsystem im Vergleich*. Berlin: Frank & Timme.
- Hacker, Martin; Ludwig, Bernd & Görz, Günther (2012). Automatisches Verstehen gesprochener Sprache. 1. Einführung in die Thematik. <https://www.yumpu.com/de/document/read/34002206/automatisches-verstehen-gesprochener-sprache-kapitel-1> (Stand: 14.07.2020)
- Hartford Business (2012). SayHi Translate app bridging the language divide. *Hartford Business Journal*. <https://www.hartfordbusiness.com/article/sayhi-translate-app-bridging-the-language-divide> (Stand: 13.07.2020)
- Havelka, Ivana (2018). *Videodolmetschen im Gesundheitswesen: Dolmetschwissenschaftliche Untersuchung eines österreichischen Pilotprojekts*. Berlin: Frank & Timme.
- Hoefert, Hans W. & Härter, Martin (Hg.) (2010). *Patientenorientierung im Krankenhaus*. Wien: Hogrefe.
- Holz-Mänttari, J. (1984). *Translatorisches Handeln: Theorie und Methode*. Helsinki: Suomalainen Tiedekatemia.
- Hutchins, W. John (1986). *Machine Translation. Past, Present, Future*. Chichester: Ellis Horwood.
- Hutchins, W. John (1995). Machine Translation. A Brief History. In: Koerner, Ernst F. K. & Asher, Ronald E. (Hg.) *Concise history of the language sciences: from the Sumerians to the cognitivists*. Oxford: Pergamon, 431-445.
- Hutchins, W. John (2015). Machine translation: history of research and applications. In: Chan, Sin-wai. (Hg.) *The Routledge Encyclopedia of Translation Technology*. London: Routledge, 120-136.
- Hutchins, W. John & Somers, H.L. (1992). *An introduction to machine translation*. London: Academic Press.

- Jüngst, Heike Elisabeth (2010). *Audiovisuelles Übersetzen: Ein Lehr- und Arbeitsbuch*. Tübingen: Narr. Francke Attempto.
- Kade, Otto (1968). *Zufall und Gesetzmäßigkeit in der Übersetzung*. Leipzig: Verlag Enzyklopädie.
- Koehn, Philipp (2010). *Statistical Machine Translation*. Cambridge: Cambridge University Press.
- Koehn, Philipp (2017). Statistical Machine Translation. Draft of Chapter 13: Neural Machine Translation. <https://arxiv.org/pdf/1709.07809.pdf> (Stand 14.07.2020).
- Kolossa, Dorothea (2017). Grundlagen der automatischen Spracherkennung. <https://www.ei.ruhr-uni-bochum.de/media/ei/lehrmaterialien/spracherkennung/19b8094b3b5cd640f4dfff0ce8afe6c978428591/SkriptASE2017d.pdf> (Stand: 13.07.2020)
- Laki, László János (2018). Mesterséges intelligencia a gépi fordításban. In: Kertész, András; Tolcsvai Nagy, Gábor & Bácskai, István. (Hg.) *A humán tudományok és a gépi intelligencia*. Budapest: Gondolat Kiadó, 156-183.
- Leanza, Y. (2007). Roles of Community Interpreters in Pediatrics as Seen by Interpreters, Physicians and Researchers. In: Pöchhacker, F. & Shlesinger, M. (Hg.) *Healthcare Interpreting: Discourse and Interaction*. Amsterdam/Philadelphia: John Benjamins, 11-34.
- Leitner, Julia (2020). *Maschinelles Dolmetschen mit iTranslate. Ein Vergleich zwischen Mensch und Maschine*. Masterarbeit, Universität Wien.
- Lionbridge (2017). Neuronale maschinelle Übersetzung: Künstliche Intelligenz erleichtert die mehrsprachige Kommunikation. <https://www.lionbridge.com/de/blog/translation-localization/neural-machine-translation-artificial-intelligence-works-multilingual-communication/> (Stand: 14.07.2020)
- Lionbridge (2018). Maschinelle Übersetzung: Wichtige Begriffe kurz erklärt. <https://www.lionbridge.com/de/blog/translation-localization/machine-translation-in-translation/> (Stand: 14.07.2020)

- Marquardt, Katja (2015). Mehrsprachigkeit in Kliniken. Dolmetschen im Krankenhaus. <https://www.goethe.de/de/kul/mol/20654851.html> (Stand: 13.07.2020)
- McGuire, Nick (2019). How accurate is Google Translate in 2019? <https://www.argotrans.com/blog/accurate-google-translate-2019/> (Stand: 13.07.2020)
- Mitrovic, Sanja (2017). *SayHi to Google Translate: Apps als GesprächsdolmetscherInnen*. Masterarbeit, Universität Wien.
- Niska, Helge (2002). Community interpreter training. Past, present, future. In: Garzone, G. & Viezzi, M. (Hg.) *Interpreting in the 21st Century*. Amsterdam/Philadelphia: John Benjamins, 133-144.
- ORF (2013). Spitäler: Ab Herbst mit Dolmetschern. <https://wien.orf.at/v2/news/stories/2599435/> (Stand: 13.07.2020)
- Panayiotou, Anita; Gardner, Anastasia; Zucchi, Emiliano; MY Goh, Anita; WH Chong, Terence; Logiudice, Dina; Lin, Xiaoping; Haralambous, Betty & Batchelor, Frances (2019). Language Translation Apps in Health Care Settings: Expert Opinion. *JMIR – Journal of Medical Internet Research* 7 (4). <https://mhealth.jmir.org/2019/4/e11316/#ref18> (Stand: 06.07.2020)
- Patil, Sumant & Davies, Patrick (2014). Use of Google Translate in medical communication: evaluation of accuracy. *BMJ – British Medical Journal* 349 (7392) <https://www.bmj.com/content/349/bmj.g7392> (Stand: 06.07.2020)
- Pfister, Beat & Kaufmann, Tobias (2017). *Sprachverarbeitung. Grundlagen und Methoden der Sprachsynthese und Spracherkennung*. Berlin/Heidelberg: Springer Verlag.
- Pöschhacker, Franz (2016). *Introducing Interpreting Studies*. London/New York: Routledge.
- Prunč, Erich (2011). Differenzierungs- und Leistungsparameter im Konferenz- und Kommunaldolmetschen. In: Kainz, Claudia; Prunč, Erich & Schögler, Rafael (Hg.) *Modelling the field of community interpreting: questions of methodology in research and training*. Wien: Lit-Verlag, 21-44.

- Ramlow, Markus (2009). *Die maschinelle Simulierbarkeit des Humanübersetzens: Evaluation von Mensch-Maschine-Interaktion und der Translatqualität der Technik*. Berlin: Frank & Timme.
- Sauerwein, Fadia (2007). Laiendolmetscher – das Zünglein an der Waage? *MDÜ – Mitteilungen für Dolmetscher und Übersetzer* 53 (5), 10-15. <http://docplayer.org/83053741-Laiendolmetscher-das-zuenglein-an-der-waage.html> (Stand 13.07.2020)
- SAVD (SAVD Videodolmetschen GmbH) <https://www.savd.at/uber-uns/> (Stand: 31.08.2020)
- SAVD Videodolmetschen (SAVD Videodolmetschen GmbH) <https://www.gfaw-thueringen.de/cms/getfile.php5?8525> (Stand: 31.08.2020)
- SayHi Translate (SayHi Amazon Company) <https://www.sayhitranslate.com> (Stand: 13.07.2020)
- Schmalz, Antonia (2019). Maschinelle Übersetzung. In: Wittpahl, Volker. (Hg.) *Künstliche Intelligenz*. Berlin/Heidelberg: Springer Vieweg, 194-211.
- Schmitz, Birte (1998). *Pragmatikbasiertes Maschinelles Dolmetschen*. Heidelberg: Groos.
- Schönert, Ulf (2012). Verständnissvolle Geräte. *Die Zeit* 2, 1–3. <https://www.zeit.de/zeit-wissen/2012/02/Spracherkennung> (Stand: 14.07.2020)
- Schulz, T. (2013). Lost in Translation. *Der Spiegel* 37, 78-79.
- Schwanke, Martina (1991). *Maschinelle Übersetzung: Ein Überblick über Theorie und Praxis*. Berlin, Heidelberg: Springer-Verlag.
- Somers, H. (2003). An overview of EBMT. In: Carl, M. & Way, A. (Hg.) *Recent Advances in Example-Based Machine Translation*. Dordrecht: Kluwer, 3-57.
- Stadt Wien | Wirtschaft, Arbeit und Statistik (o.J.). Bevölkerung nach Staatsangehörigkeit und Geschlecht 2018 und 2019.

<https://www.wien.gv.at/statistik/bevoelkerung/tabellen/bevoelkerung-staat-geschl-zr.html>

(Stand: 14.07.2020)

Stein, Daniel (2009). Maschinelle Übersetzung – ein Überblick. *JLCL – Journal for Language Technology and Computational Linguistics* 24 (3), 5-18.

Straub, Michael (2016). Brauchen Ärzte und Spitäler einen Dolmetscher? https://www.mplaw.at/media/dokumente/2016-03-31_oekz-3-4-2016_brauchen-aerzte-und-spitaeler-einen-dolmetscher_mis_print_1497.pdf (Stand: 14.07.2020)

Tripathi, S. & Sarkhel, J.K. (2010). Approaches to machine translation. *Annals of Library and Information Studies* 57 (1), 388-393.

Voicepoint AG (o.J.). Spracherkennung. <https://voicepoint.ch/angebot/spracherkennung/> (Stand: 14.07.2020)

Wahlster, Wolfgang (2000). Mobile Speech-to-Speech Translation of Spontaneous Dialogs: An Overview of the Final Verbmobil System. In: Wahlster, Wolfgang. (Hg.) *Verbmobil: Foundations of Speech-to-Speech Translation*. Berlin: Springer-Verlag, 3-21.

Waibel, Alexander (2015). Sprachbarrieren durchbrechen: Traum oder Wirklichkeit? *Nova Acta Leopoldina* 122 (410), 101-123.

Waibel, Alexander & Fügen, Christian (2008). Spoken Language Translation. *IEEE Signal Processing Magazine* 25 (3), 70-79.

Waibel, A.; Jain, A.; McNair, A.; Saito, H.; Hauptmann, A. & Tebelskis, J. (1991). *JANUS: A speech-to-speech translation system using connectionist and symbolic processing strategies. Proceedings of the International Conference on Acoustics, Speech and Signal Processing*. Toronto: ICASSP.

Werthmann, Antonia & Witt, Andreas (2014). Maschinelle Übersetzung – Gegenwart und Perspektiven. In: Stickel, Gerhard. (Hg.) *Translation and Interpretation in Europe*. Frankfurt am Main: Lang, 79-103.

Wetz, Andreas (2016). In Wien wohnen Menschen aus 183 Ländern. *Die Presse*. <https://die-presse.com/home/panorama/wien/5107501/In-Wien-wohnen-Menschen-aus-183-Laendern-2016> (Stand 13.07.2020)

Anhang

Transkription des Experimentes

(Experiment durchgeführt am 07.12.2019)

	Gesprochene Eingabe	ASR + Transkription	MÜ	Bildbeschreibung, nonverbale Kommunikation
	A	B	C	D
1	Hallo!	x	x	<i>gibt die Hand</i>
2	Bitte, setzen Sie sich!	x	x	<i>schüttelt die Hand und zeigt auf den Stuhl</i>
3	<u>Jó</u> napot kívánok doktor úr, köszönöm szépen, hogy időt szánt rám <i>pont</i> . Bécsben dolgozom, de sajnos nem beszélek németül, ezért hoztam el ezt az applikációt, ami remélem segít nekünk a kommunikációban <i>pont</i> .	Napot kívánok doktor úr köszönöm szépen hogy időt szánt rám. Bécsben dolgozom, de sajnos nem beszélek németül ezért hoztam el ezt az applikációt amit remélem segít nekünk majd a kommunikációban.	Ich wünsche Ihnen einen Tag, Doktor. Ich arbeite in Wien, aber leider spreche ich kein Deutsch, deshalb habe ich diese App mitgebracht, die uns hoffentlich helfen wird, zu kommunizieren.	<i>setzt sich an den Schreibtisch und beginnt zu sprechen; der Arzt möchte gleich antworten und greift nach dem Tablet, der Patient erhebt jedoch den Zeigefinger, um zu signalisieren, dass er noch etwas sagen möchte</i>
4	Gyorsan elmagyarázom, hogy hogyan működik a program. Itt vannak a nyelvek, amikor beszélünk, ezeket a gombokat kell megnyomni és utána kell beszélni <i>pont</i> .	Gyorsan elmagyarázom hogy hogyan működik a program, itt vannak a nyelvek, amikor beszélünk ezeket a gombokat kell megnyomni és utána kell beszélni.	Ich erkläre schnell, wie das Programm funktioniert, hier sind die Sprachen, wenn wir über diese Tasten sprechen, die gedrückt werden und dann sprechen.	<i>zeigt während der Sprach-eingabe auf die Tasten; nachdem das Gerät übersetzt hat dreht der Arzt das Tablet zu sich</i>
5	Gut, ich glaube, das habe ich verstanden <i>Punkt</i> . Darf ich Sie fragen, was ich für Sie tun kann?	Ich glaube das habe ich verstanden. Darf ich Sie fragen was ich für Sie tun kann	Azt hiszem, megértettem. Megkérdezhetem, hogy mit tehetek érted	

6	Visszértágulataim vannak enyhe fokban a lábamon és ez egy... ez ügyben szeretnék tanácsot kérni <i>pont.</i>	Visszértágulataim vannak enyhe fokban a lábamon és ez egy ez ügyben szeretnék tanácsot kérni.	Ich habe Krampfadern in einem leichten Grad auf meinem Bein, und das ist eine, die ich gerne um Rat fragen würde.	
7	Darf ich fragen, ist es das rechte Bein oder das linke Bein <i>Punkt</i> , und haben Sie schon irgendeine Form von Behandlung deswegen gehabt?	Darf ich fragen, ist es das rechte Bein oder das linke Bein. Und haben Sie schon irgendeine Form von Behandlung deswegen gehabt	Megkérdezhetem, ez a jobb láb vagy a bal lábát. És volt már bármilyen formában a kezelés e	
8	<u>Mindkét</u> lábamon tapasztalom <i>pont.</i> De csak enyhe fokban <i>pont.</i>	<u>Mind a két</u> lábamon tapasztalom. De csak enyhe fokban.	Ich erlebe es auf beiden Beinen. Aber nur in geringem Maße.	
9	Haben Sie einen Beruf, bei dem Sie viel stehen oder sitzen müssen <i>Komma</i> , tut es im Sommer mehr weh als im Winter?	Haben Sie einen Beruf bei dem <u>sie</u> viel stehen oder sitzen müssen, tut es im Sommer mehr weh als im Winter	Ha van egy munka, ahol meg kell állni vagy ülni sokat, hogy jobban fáj a nyáron, mint télen	
10	Télen vagy nyáron, ezt nem figyeltem meg, mindegy <i>pont.</i> Ülőmunkát végzek általában, viszonylag sokat ülök <i>pont.</i> Ami aggaszt, ami gondot okoz, hogy a családban az édesapámnál és a nagyszüleimnél is voltak visszértágulati problémák, ezért aggódok egy kicsit <i>pont.</i>	Télen vagy nyáron ezt nem figyeltem meg mindegy. Ülőmunkát végzek, általában viszonylag sokat ülök. Ami aggaszt, ami gondot okoz, hogy a családban az édesapámnál és a nagyszüleimnél is voltak visszértágulati problémák ezért aggódok egy kicsit.	Im Winter oder Sommer habe ich das sowieso nicht gemerkt. Ich mache einen Sit-down-Job, normalerweise siss eitere ich relativ viel. Was mich beunruhigt, was ein Problem ist, ist, dass mein Vater und meine Großeltern Krampfadern in der Familie hatten, also mache ich mir ein wenig Sorgen.	

11	<u>Spüren</u> Sie eine unangenehme... unangenehmen <u>Schmerz in den</u> Beinen wenn sie länger sitzen oder gibt es <u>auch</u> Schwellungen der Beine <u>Punkt</u>? Wenn Sie Schwellungen haben, ist es <u>eh</u>er die <u>Innenseite</u> oder die Außenseite?	<u>Spielen</u> Sie eine unangenehme unangenehmen <u>schmerzenden</u> Beinen wenn sie länger sitzen oder gibt es <u>noch</u> Schwellungen der Beine. Wenn Sie Schwellungen haben ist es <u>hier</u> die <u>ine</u> Seite oder die Außenseite	Játsszon kellemetlen kellemetlen lábakkal, amikor hosszabb ideig ülnek, vagy még a lábak duzzanata is fennáll. Ha van duzzanat itt ez a belső oldalán, vagy a külső	
12	Elnézést, de ezt <u>nem értettem</u> pontosan. Ismétélje meg, kérem <u>pont</u>.	Elnézést, de ezt <u>nemértettem</u> pontosan ismétélje meg kérem.	Es tut mir leid, aber ich habe es nicht verstanden, bitte wiederholen Sie es genau, bitte.	
13	Ich frage genauer nach, welche Beschwerden Ihnen die Venen in den Beinen machen <u>Punkt</u>. Haben <u>Sie</u> Schmerzen, wenn <u>Sie</u> sich wenig bewegen und viel sitzen <u>Punkt</u>? Haben Sie Schwellungen in den Knöchelregionen der Beine <u>Punkt</u>? Wenn Sie Schwellungen haben <u>Komma</u>, sind es <u>eh</u>er auf der Innenseite des Beines oder auf der Außenseite <u>des Beines</u>?	Ich frage genauer nach, welche Beschwerden Ihnen die Venen in den Beinen machen. Haben <u>sie</u> Schmerzen wenn <u>sie</u> sich wenig bewegen und viel sitzen. Haben Sie Schwellungen in den Knöchelregion der Beine. Wenn Sie Schwellungen haben, sind es <u>dir</u> auf der Innenseite des Beines oder auf der Außenseite	Kérem, többet arról, milyen kellemetlen a vénák a lábad neked. Fájdalma van, amikor keveset mozognak és sokat ülnek. Van duzzanat a boka régióban a lábak. Ha van duzzanat, ez akkor a belsejében a láb vagy a külső	
14	Általában a vádlim feszes, kemény, ezt érzem gyakran. Bár komoly problémát nem okoz, igazán azt szeretném kérdezni, hogyan tudnám a komolyabb bajokat megelőzni <u>pont</u>.	Általában a vádlim feszes, kemény ezt érzem gyakran. Bár komoly problémát nem okoz igazán azt szeretném kérdezni, hogyan tudnám a komolyabb bajokat megelőzni.	Normalerweise ist meine Wade eng, so oft schwer zu fühlen. Obwohl es nicht wirklich ein ernstes Problem ist, möchte ich Sie fragen, wie ich ernstere Probleme verhindern kann.	

15	Bevor ich Ihnen das <u>beantworte</u>, möchte ich Sie gerne untersuchen <u>Punkt</u>. Dazu werde ich mir <u>Ihre</u> Beine im Stehen anschauen und einen Ultraschall machen <u>Punkt</u>. Dazu müssen Sie bitte in der Garderobe hinter <u>Ihnen</u> die Beine freimachen <u>Punkt</u>. Bitte ziehen Sie Ihre Hose, die Schuhe und die Strümpfe aus und stellen Sie sich auf das kleine <u>Podest Komma</u>, wo die Füße grün eingezeichnet sind.	Bevor ich Ihnen das <u>beantwortet</u> möchte ich Sie gerne untersuchen. Dazu werde ich mir <u>ihre</u> Beine im Stehen anschauen und einen Ultraschall machen. Dazu müssen Sie bitte in der Garderobe hinter <u>ihnen</u> die Beine freimachen. Bitte ziehen Sie Ihre Hose die Schuhe und die Strümpfe aus und stellen Sie sich auf das kleine <u>Protest</u>, wo die Füße grün eingezeichnet sind	Mielőtt válaszolnék nektek, hogy szeretném megvizsgálni téged. Fogom nézni a lábát állva, és nem egy ultrahang. Ehhez kérjük, hogy törölje a lábát az öltözőben mögöttük. Kérjük, vegye le a nadrágját, és harisnya és állni a kis tiltakozás, ahol a lábad készült zöld	
16	Jó, rendben.	Jó rendben	Alles in Ordnung, in Ordnung.	<i>Der Arzt zeigt dem Patienten, wo die Garderobe ist, der Patient macht sich die Beine frei und stellt sich auf das Podest. Der Arzt beginnt mit der Untersuchung</i>
17	Bitte zeigen Sie mir an <u>Ihren</u> Beinen die Bereiche, in denen <u>Sie</u> Beschwerden spüren, wenn <u>Sie</u> länger sitzen.	Bitte zeigen Sie mir an <u>ihren</u> Beinen die Bereiche in denen <u>sie</u> Beschwerden spüren, wenn <u>sie</u> länger sitzen	Kérjük, mutassa meg a lábukat a területeken, ahol úgy érzik, kényelmetlenséget, ha ülünk tovább	<i>der Patient zeigt auf seine Waden</i>
18	Ok. Bitte bleiben Sie gleich <u>so stehen</u>, ich werde jetzt an mehreren Stellen <u>eine</u> Ultraschallsonde hinhalten und dabei gleichzeitig auf <u>Ihre</u> Waden drücken <u>Punkt</u>. Es wird Ihnen nichts dabei <u>wehtun</u>.	O. K. Bitte bleiben Sie gleich <u>zustehen</u>, ich werde jetzt an mehreren Stellen <u>in der</u> Ultraschallsonde hinhalten und dabei gleichzeitig auf <u>ihre</u> Waden drücken. Es wird Ihnen nichts dabei <u>wehtuend</u>	Ok. Kérem, maradjon rendben, én most megáll több helyen az ultrahangos szondán, és nyomjuk meg a borjak ugyanabban az időben. Semmi sem fog bántani	<i>der Arzt möchte mit der Untersuchung anfangen, muss aber auf die Übersetzung warten, da aber die App hier ein wenig länger braucht, zeigt er ungeduldig auf den Tablet</i>

19	Bitte atmen Sie tief ein und <u>husten</u> Sie, wenn ich jetzt sage.	Bitte atmen Sie tief ein und <u>wussten</u> Sie, wenn ich jetzt sage	Kérjük, hogy egy mély lélegzetet, és tudom, mikor azt mondom, most	<i>Patient atmet tief ein</i>
20	Jetzt!	x	x	
21				<i>atmet tief ein</i>
22				<i>hustet, um zu zeigen, was er gemeint hat</i>
23				<i>hustet nach</i>
24	Ok, Stopp.	x	x	
25	Jetzt!	x	x	
26				<i>hustet</i>
27	Bitte <u>stellen</u> Sie jetzt das linke Bein mit einem halben Schritt Weg nach vorne weg <i>Punkt.</i>	<u>Wählen</u> Sie jetzt das linke Bein mit einem halben Schritt Weg nach vorne weg.	Most válassza ki a bal lábát egy fél lépéssel.	<i>Während der Arzt auf die Übersetzung wartet, macht er kreisende, ungeduldige Bewegungen mit den Händen. Er schiebt das linke Bein des Patienten einen halben Schritt nach vorne.</i>
28				<i>hustet</i>
29				<i>hustet nach</i>
30	<u>Stellen Sie bitte</u> jetzt das linke Bein zu mir nach hinten und das rechte einen halben Schritt nach vorne.	Jetzt das linke Bein zu mir nach hinten und das Rechte einen halben Schritt nach vorne	Most a bal láb nekem hátra, és a jobb fél lépés előre	
31	Noch mal husten	x	x	<i>hustet</i>
				<i>hustet nach</i>
32	Drehen Sie sich jetzt bitte zu mir nach vorne und <u>halten die Hose</u> ein wenig in den Schritt hinein <i>Komma</i> , damit ich bei der Leiste messen kann.	Drehen Sie sich jetzt bitte zu mir nach vorne und <u>halt in die Hose</u> ein wenig in den Schritt hinein, damit ich bei der Leiste messen kann	Most kérem, forduljon hozzám előre, és tartsa a nadrágot egy kicsit az y, hogy tudom mérni a Lágyulás	
33				<i>hustet</i>
				<i>hustet nach</i>

34	Danke, mit der Untersuchung sind wir fertig <i>Komma</i> , Sie können sich wieder anziehen und setzen sich zu mir an den Schreibtisch.	Danke mit der Untersuchung sind wir fertig, Sie können sich wieder anziehen und setzen sich zu mir an den Schreibtisch	Köszönöm a vizsgálatot végeztünk, akkor újra felöltözni, és üljön le velem az íróasztalon	<i>Der Patient nickt, der Arzt streckt den Daumen nach oben</i>
35	Ich werde Ihnen nun anhand einer Grafik ihr Problem an den Beinvenen erklären.	Ich werde Ihnen nun anhand einer Grafik ihr Problem an den Beinvenen erklären	Én most használ egy grafikus megmagyarázni a problémát a láb vénák	<i>Arzt und Patient setzen sich an den Schreibtisch, Der Arzt erklärt das Problem anhand einer Grafik. Er verwendet einen Kugelschreiber, um die erwähnten Venen und Mechanismen besser zeigen zu können.</i>
36	Diese Zeichnung ist eine vereinfachte Darstellung der Venen des Beines <i>Punkt</i> . Man muss sich die blauen Linien räumlich als Röhren vorstellen <i>Punkt</i> .	Diese Zeichnung ist eine vereinfachte Darstellung der Venen des Beines. Man muss sich die blauen Linien räumlich als Röhren vorstellen.	Ez a rajz egy egyszerűsített ábrázolása a vénák a lábát. Meg kell gondolni, a kék vonalak a csöveket.	
37	Es gibt eine große <u>Vene</u> ganz in der Tiefe der Muskulatur, die das meiste Blut sammelt und direkt zum Herzen führt <i>Punkt</i> .	Gibt eine große <u>Biene</u> ganz in der Tiefe der Muskulatur, die das meiste Blut sammelt und direkt zum Herzen führen.	Ad egy nagy méh a mélység az izmok, hogy összegyűjti a legtöbb vért, és vezet közvetlenül a szívbe.	
38	Diese tiefe Beinvene wird unterstützt durch zwei dünne <u>Venen</u> , die unter der Haut laufen und das Blut von der <u>Hautregion</u> einsammeln.	Diese tiefe Beinvene wird unterstützt durch zwei dünne <u>Bienen</u> die unter der Haut laufen und das Blut von der <u>Haut Region</u> einsammeln	Ez a mély láb véna támasztja alá két vékony méhek fut a bőr alatt, és összegyűjti a vért a bőr régió	

39	Es ist nicht selbstverständlich, dass das Blut gegen die Schwerkraft hinauf zum Herzen fließt.	Es ist nicht selbstverständlich dass das Blut gegen die Schwerkraft hinauf zum Herzen fließt	Ez nem magától értetődő, hogy a vér áramlik fel a szívbe a gravitáció ellen	
40	Dazu braucht es zwei Mechanismen <i>Punkt.</i>	Dazu braucht es zwei Mechanismen.	Ehhez két mechanizmusra van szükség.	
41	<u>Durch Bewegung</u> der Muskulatur kommt es zu einem Druck auf die <u>Venen</u> , dieser Druck bringt das Blut zu fließen.	<u>Durchbewegung</u> der Muskulatur kommt es zu einem Druck auf die <u>Wiener</u> dieser Druck bringt das Blut zufließen	Az izmok mozgása van a nyomás a bécsi ez a nyomás hozza a vér áramlását	
42	Damit das Blut aber nur in die Richtung "Herz" und von der <u>Außenvene</u> zur inneren <u>Vene</u> fließen kann, braucht es eine Venenklappen, die das Blut nur in die richtige Richtung fließen lassen.	Damit das Blut aber nur in die Richtung Herz und von der <u>Außenlinie</u> zur inneren <u>Biene</u> fließen kann, braucht es eine Venenklappen die das Blut nur in die richtige Richtung fließen lassen	De a vér áramlását csak az irányba a szív és a külső vonal a belső méh, akkor szükség van a véna szelep, amely csak lehetővé teszi a vér áramlását a helyes irányba.	
43	Bei vielen Patienten bestehen <u>Venenprobleme</u> , weil eine Klappe nicht funktioniert <i>Komma</i> , das Blut umdreht und wieder <u>Richtung</u> Beine fließt, wie ich es <u>Ihnen</u> kurz auf der Zeichnung zeigen werde.	Bei vielen Patienten bestehen <u>Bienen Probleme</u> , weil eine Klappe nicht funktioniert, das Blut umdreht und wieder <u>Richtungs</u> Beine fließt wie ich es <u>ihnen</u> kurz auf der Zeichnung zeigen werde	Sok beteg problémái vannak a méhek, mert a fedél nem működik, a vér fordul, és a vér áramlik vissza a lábak felé, mint én megmutatom nekik röviden a rajz	

44	Das ist bei <u>Ihnen</u> aber nicht der Fall <i>Komma</i> , die Klappen funktionieren alle <u>sehr gut</u> .	Das ist bei <u>ihnen</u> aber nicht der Fall, die Klappen funktionieren alle <u>sercan</u>	De nem ez a helyzet velük, a fűleket minden munkát sercan	
45	Die Klappen funktionieren bei Ihnen alle sehr gut	Die Klappen funktionieren bei Ihnen alle sehr gut	A fűleket nagyon jól működik az Ön számára	
46	Damit haben Sie das Problem <i>Komma</i> , dass ihr Vater und ihr Großvater offenbar gehabt haben, nicht <u>geerbt</u> .	Damit haben Sie das Problem, dass ihr Vater und ihr Großvater offenbar gehabt haben nicht <u>geärgert</u>	Ez azt jelenti, hogy a probléma, hogy apja és nagyapja láthatóan nem bosszantott	
47	Nicht <u>geerbt</u> .	Nicht <u>geärgert</u>	Nem bosszús	<i>erkennt, dass das Ende des Satzes nicht richtig transkribiert wurde und wiederholt die letzten zwei Wörter</i>
48				<i>zeigt "es spinnt"</i>
49		Damit haben Sie das Problem, dass ihr Vater und ihr Großvater offenbar gehabt haben nicht <u>geerbt</u>	Ez azt jelenti, hogy nem örökölte a problémát, hogy apja és nagyapja láthatóan	<i>Der Patient hilft dem Arzt, die Korrekturfunktion zu verwenden; der Arzt korrigiert den Satz mit Hilfe der Tastatur (geärgert - geerbt)</i>
50	Wenn sie aber viel sitzen müssen, fällt die Komponente der Muskelpumpe weg.	Wenn sie aber viel sitzen müssen, fällt die Komponente der Muskelpumpe weg	Azonban, ha kell ülni sokat, az alkatrész az izom szivattyú elesik	

51	Dieses Problem können Sie nur lösen, indem sie während <u>Ihrer</u> Berufstätigkeit immer wieder versuchen <i>Komma</i>, aufzustehen und kurze Wege zu Fuß zu erledigen <i>Punkt</i>. Weiters sollten Sie in der Freizeit sich bemühen, viel Bewegung mit den Beinen zu machen <i>Punkt</i>. Auch kann ich Ihnen ein Medikament verschreiben, dass dieses Problem mildern wird.	Dieses Problem können Sie nur lösen, indem sie während <u>ihrer</u> Berufstätigkeit immer wieder versuchen, aufzustehen und kurze Wege zu Fuß zu erledigen. Weiters sollten Sie in der Freizeit sich bemühen, viel Bewegung mit den Beinen zu machen. Auch kann ich Ihnen ein Medikament verschreiben, dass dieses Problem mildern wird	Csak akkor tudja megoldani ezt a problémát, ha többször próbál feljutni, és gyakran rövid távolságok munka közben. Azonkívül, Önnek kellene csinál egy erőfeszítésben-a szabad idő-hoz csinál sok gyakorlással-a lábak. Is tudok felírni neked egy gyógyszer, amely enyhíteni ezt a problémát	
52	Haben Sie zu dieser Erklärung irgendwelche Fragen?	Haben Sie zu dieser Erklärung irgendwelche Fragen	Kérdése van ezzel a nyilatkozattal kapcsolatban?	
53				<i>schüttelt den Kopf</i>
54	<u>Köszönöm szépen</u>, mindent megértettem. Igyekszem többet mozogni és az életmódomat is úgy alakítani, hogy minél kevesebb rizikó legyen benne.	<u>Köszönömszépen</u> mindent megértettem igyekszem többet mozogni és az életmódomat is úgy alakítani, hogy minél kevesebb Rizikó legyen benne	Vielen Dank Für alles, was ich verstehe, versuche ich, mich mehr zu bewegen und meinen Lebensstil so zu gestalten, dass es so wenig Risiko wie möglich hat.	
55	Das ist sehr gut <i>Punkt</i>. Ich schreibe Ihnen das Medikament auf <i>Komma</i>, es ist ein Extrakt aus Pflanzen und damit ohne Nebenwirkungen <i>Punkt</i>.	Das ist sehr gut. Ich schreibe Ihnen das Medikament auf, es ist ein Extrakt aus Pflanzen und damit ohne Nebenwirkungen.	Ez nagyon jó. Írom le a gyógyszert neked, ez egy kivonat a növények, és ezért mellékhatások nélkül.	

56	Man nimmt es zweimal täglich als Tablette ein <i>Komma</i> , und dies genügt für eine Dauer von 30 Tagen <i>Punkt.</i>	Man nimmt es zweimal täglich als Tablette ein, und dies genügt für eine Dauer von 30 Tagen.	Meg kell venni naponta kétszer, mint egy tablettát, és ez elegendő egy 30 napos időtartamra.	
57	Noch eine Frage <i>Komma</i> , betreiben Sie <u>Sport</u> , in der Natur oder einem Fitnessstudio?	Noch eine Frage, betreiben Sie <u>spart</u> , in der Natur oder einem Fitnessstudio	Egy másik kérdés, hogy mentse a természetben vagy egy tornaterem	
58	Nemrégem <u>kezdtm</u> el ööö... tornaterembe, kondicionáló terembe jární, remélem az is segíteni fog.	Nemrégem <u>keztem</u> el <u>tudom</u> tornaterembe kondicionáló terembe jární remélem az is segíteni fog.	Ich habe in der Lage, in ein Fitnessstudio vor langer Zeit zu gehen, so hoffe ich, es hilft.	
59	Wenn Sie ins Fitnessstudio gehen, Sport betreiben oder sich körperlich anstrengen, kommt jetzt noch von mir ein wichtiger Tipp <i>Komma</i> , um ihre <u>Venen</u> auch in Zukunft gesund zu halten.	Wenn Sie ins Fitnessstudio gehen, Sport betreiben oder sich körperlich anstrengen kommt jetzt noch von mir ein wichtiger Tipp, um ihre <u>Bienen</u> auch in Zukunft gesund zu halten	Ha megy a tornaterem, a testmozgás és a testmozgás kemény, adok egy fontos tipp, hogy tartsa a méhek egészségesek a jövőben	
60	Ich habe <u>Sie</u> bei der Untersuchung mehrfach aufgefordert, zu <u>husten</u> <i>Punkt.</i> <u>Husten</u> bedeutet eine Kontraktion mit dem Zwerchfell, dies ist der größte und stärkste Muskel in unserem Körper.	Ich habe <u>sie</u> bei der Untersuchung mehrfach aufgefordert zu <u>wussten</u> . <u>Houston</u> bedeutet eine Kontraktion mit dem Zwerchfell, dies ist der größte und stärkste Muskel in unserem Körper	Kérdeztem tőle több idő alatt a vizsgálat-hoz tud. Houston azt jelenti, összehúzódása a rekeszizom, ez a legnagyobb és legerősebb izom testünk.	

61	Ich habe <u>Sie</u> bei der Untersuchung mehrfach aufgefordert, zu <u>husten</u> <i>Punkt</i> . <u>Husten</u> bedeutet eine Kontraktion mit dem Zwerchfell, dies ist der größte und stärkste Muskel in unserem Körper.	Ich habe <u>sie</u> bei der Untersuchung mehrfach aufgefordert zu <u>husten</u> . <u>Husten</u> bedeutet eine Kontraktion mit dem Zwerchfell, dies ist der größte und stärkste Muskel in unserem Körper	Megkértem őt, hogy köhögést többször a vizsgálat során. Köhögés: összehúzódása a rekeszizom, ez a legnagyobb és legerősebb izom testünk	<i>verwendet die manuelle Korrekturfunktion, um die falsch erkannten Wörter (wussten - husten, Houston - Husten) zu korrigieren</i>
62	Na!	x	x	<i>ärgert sich</i>
63	Na, jetzt geht's.	x	x	
64	Beim Husten oder <u>Pressen</u> mit dem Bauch wird mit großem Druck das Venenblut wieder in die Beine zurückgedrückt.	Beim Husten oder <u>pressend</u> mit dem Bauch wird mit großem Druck das Venenblut wieder in die Beine zurückgedrückt	Amikor köhöt vagy présel a hasüregbe, a véna vére nagy nyomással visszaterekedik a lábakra	
65	Dies belastet die Venenklappen und kann zu Schäden führen, wie sie in der Bezeichnung, die ich jetzt aufgedeckt habe, gezeichnet sind.	Dies belastet die Venenklappen und kann zu Schäden führen wie sie in der Bezeichnung, die ich jetzt aufgedeckt habe gezeichnet sind	Ez teszi a megterhelést a véna fűleket és károsíthatja a károkat, mivel azok szerepelnek a kijelölésben már fedeztetlen	
66	Damit die Venen nicht krank werden <i>Komma</i> , achten Sie bitte bei sportlicher Belastung und körperlicher Anstrengung immer darauf, dass sie nicht mit dem Bauch pressen <i>Punkt</i> . Dies tut man üblicherweise, um sich die Belastung zu erleichtern <i>Punkt</i> . Es macht aber einen Druck auf die Venenklappen <i>Komma</i> , und kann eventuell auch noch später zu Problemen mit den <u>Venen</u> führen <i>Punkt</i> .	Damit die Venen nicht krank werden, achten Sie bitte bei sportlicher Belastung und körperliche Anstrengung immer darauf, dass sie nicht mit dem Bauchpressen. Dies tut man üblicherweise, um sich die Belastung zu erleichtern. Es macht aber einen Druck auf die Venenklappen, und kann eventuell auch noch später zu Problemen mit den <u>Bienen</u> führen.	Annak megakadályozására, hogy a vénák szerzés beteg, kérjük, mindig győződjön meg arról, hogy nem nyomja a has, ha gyakorlása és a kifejtése. Ez általában történik, hogy enyhítse a törzs. Ez azonban nyomást gyakorol a véna flaps, és problémákat okozhat a méhek még később.	

67	Damit die Venen nicht krank werden, achten Sie bitte bei sportlicher Belastung und körperlicher Anstrengung immer darauf, dass sie nicht mit dem Bauch pressen. Dies tut man üblicherweise, um sich die Belastung zu erleichtern. Es macht aber einen Druck auf die Venenklappen, und kann eventuell auch noch später zu Problemen mit den <u>Venen</u> führen.	Damit die Venen nicht krank werden, achten Sie bitte bei sportlicher Belastung und körperlicher Anstrengung immer darauf, dass sie nicht mit dem Bauch pressen. Dies tut man üblicherweise, um sich die Belastung zu erleichtern. Es macht aber einen Druck auf die Venenklappen, und kann eventuell auch noch später zu Problemen mit den <u>Venen</u> führen.	Annak megakadályozására, hogy a vénák szerzés beteg, kérjük, mindig győződjön meg arról, hogy nem nyomja a has, ha gyakorlása és a kifejtése. Ez általában történik, hogy enyhítse a törzs. Ez azonban nyomást gyakorol a véna flaps, és problémákat okozhat a vénák még később.	
68	Nagyon szépen köszönöm a vizsgálatot és a hasznos tanácsot. A gyógyszert be fogom venni, illetve ezt a <u>gyógyszert meg fogom</u> vásárolni és remélem hogy	Nagyon szépen köszönöm a vizsgálatot és a hasznos tanácsot a gyógyszert be fogom venni, illetve ezt a <u>gyógyszert megfogom</u> vásárolni és remélem hogy	Vielen Dank für die Untersuchung und hilfreiche Ratschläge, die ich nehmen und kaufen dieses Medikament und hoffen, dass	
69	Remélem, hogy hasonló jó egészségben találkozunk a jövőben <i>pont</i>.	Remélem, hogy hasonló jó egészségben találkozunk a jövőben.	Ich hoffe, Sie in Zukunft in ähnlicher Gesundheit zu sehen.	
70	Ich wünsche Ihnen eine gute Heimfahrt, <i>Komma</i> nett, dass <u>Sie</u> bei mir waren <i>Rufzeichen</i>!	Ich wünsche Ihnen eine gute Heimfahrt, nett dass <u>sie</u> bei mir waren!	Szeretném, ha egy jó utazás haza, szép, hogy voltak velem!	

Interview mit dem Arzt (IA)

Durchgeführt am 07.12.2019

1. Wie fanden Sie die ganze Situation? Wie fanden Sie die App, war sie einfach zu verwenden, war sie logisch aufgebaut?

Also im Prinzip ist die App logisch aufgebaut, einfach zu verwenden und gut. Wenn man aber in der Situation ist, dass man miteinander spricht, wäre es natürlich schlau, wenn die App – wie ein Tennisplatz – eine Mitte hat und sie schreibt einmal für mich richtig und einmal für den Gesprächspartner richtig, weil das kostet Zeit, das Hin-und Herdrehen.

2. Ja, genau. Haben Sie die manuelle Korrekturfunktion verwendet?

Ja.

3. War es einfach zu verwenden, oder war es eher störend, immer wieder schreiben zu müssen?

Es ist dann störend, wenn das Gerät mich nicht gut versteht, aber das Gerät versteht mich meistens nicht gut, wenn ich in dem Inhalt, in dem, was ich sagen möchte, unsicher bin und zaudere. Das merkt dieser Apparat sehr schnell.

4. Waren diese Pausen, das Warten störend?

Es könnte natürlich schneller gehen, üblicherweise würde ich für den Patienten eine Viertelstunde brauchen, so brauche ich doppelt so lang.

5. Hatten Sie den Eindruck, dass der Patient Sie gut verstanden hat, oder hat er nur so getan, als ob er alles gut verstanden hätte?

Also ich glaube schon, dass er alles gut verstanden hat.

6. Falls jemand in der Zukunft mit dieser App kommen würde, wie würden Sie reagieren? Würden Sie auf einen Dolmetscher bestehen, oder es so probieren?

Ich würde mir mehr Zeit einteilen und würde es probieren mit der App.

7. Vielen Dank für die Antworten. Gibt es vielleicht etwas, was Sie zu der ganzen Situation hinzufügen würden?

Vielleicht nur, dass ich im Spital...wir haben sehr viele Sachen, die wir diktieren müssen, also wir machen den Operationsbericht, wir machen Arztbriefe, das wird alles über so...elektronische Apps gemacht und ich habe daher gelernt, deutlich zu sprechen, oder sag ma's so... deutlicher zu sprechen, damit der Apparat einen gut verstehen kann. Ich kann mir aber gut vorstellen, dass viele das nicht können und dass dann die Anwendbarkeit der App damit mehr limitiert ist, weil natürlich das Korrigieren, das nicht genaue Verstehen des Apparats sinnverzerrend wirkt und dann ist die Kommunikation gestört.

Interview mit dem Patienten

Durchgeführt am 09.12.2019

(Übersetzt aus dem Ungarischen)

1. Wie lange hat es gedauert, sich in der App zurechtzufinden, als Sie sie zum ersten Mal verwendet haben? Denken Sie, dass SayHi Translate logisch und transparent aufgebaut ist?

Meiner Meinung nach ist die App sehr einfach und logisch aufgebaut. Ich denke, dass ich ein bisschen besser mit Mobilien Apps umgehe als der Durchschnitt meiner Altersklasse, da ich mein Smartphone und Tablet tagtäglich benutze. Deswegen konnte ich die Verwendung der App relativ schnell und einfach lernen, nach ungefähr 15 Minuten Übung habe ich mich gut zurechtgefunden.

2. Während der Untersuchung wurden die Wörter des Arztes auf einer weiblichen Stimme von der App „gedolmetscht“. Fanden Sie dies störend?

Die Tatsache, dass der Arzt ein Mann war, die Dolmetsch-App aber auf einer weiblichen Stimme gesprochen hat, war gar nicht störend. Das ist auch bei einem „traditionellen“ Dolmetscheinsatz normal.

3. Zwischen der Spracheingabe und der „Dolmetschung“ brauchte die App immer paar Sekunden, die angekommenen Informationen zu verarbeiten. War dieses Warten unangenehm oder störend?

Die Textverarbeitung der App führte tatsächlich oft zum „Leerlauf“ in der Kommunikation. Meiner Ansicht nach war dies ein bisschen störend, da das Gespräch so immer abgebrochen wurde, ich finde die App aber grundsätzlich verwendbar.

4. Die ungarische Übersetzung war sehr oft sowohl grammatikalisch als auch syntaktisch gesehen fehlerhaft, trotzdem konnte die Kommunikation im Großen und Ganzen zustande kommen. War es schwer oder anstrengend, den Sinn dieser fehlerhaften Sätze zu verstehen? Gab es Stellen, wo Sie sogar den Faden völlig verloren haben?

Na ja... das Warten auf die Übersetzung war für mich nicht störend, wegen der Fehlübersetzungen habe ich aber leider mehrmals den Faden verloren, da das Gerät nicht einmal annähernd das übersetzt hat, worüber wir gesprochen haben. Hätte ich ernstere Beschwerden gehabt, hätte ich wahrscheinlich nicht verstanden, was der Arzt mir mitteilen wollte. Oder vielleicht hätte er mich nicht verstanden.

5. Gab es während des Gespräches Stellen, wo Sie den Arzt nicht verstanden haben? Wenn ja, haben Sie zurückgefragt?

Ja, wie gesagt, gab es Stellen, wo ich den Arzt tatsächlich nicht verstanden habe. Da ich aber schon ein wenig Erfahrung mit der App hatte, wusste ich, dass sie das Gesagte nicht immer gut versteht und manchmal Fehlübersetzungen liefert. Der Arzt hat sehr langsam gesprochen, ich wusste, dass nicht er für die schlechten Übersetzungen verantwortlich ist, sondern die Applikation selbst. Er hätte keinen Sinn gehabt, nachzufragen. Vor allem weil wir nicht

z.B. einen wichtigen Operationstermin vereinbart haben, er hat „lediglich“ Tipps für eine gesunde Lebensführung gegeben.

6. Haben Sie während der Untersuchung die manuelle Korrekturfunktion verwendet? War sie einfach zu verwenden? Inwieweit fanden Sie es störend, dass die Kommunikation im Grunde genommen mittels eines Tablets gelaufen ist?

Die manuelle Korrektur war gar nicht störend, sie machte aber die ärztliche Untersuchung irgendwie wirklichkeitsfremd. Ich würde das Tablet bzw. diese App nicht gerne in einer solchen Situation in Anspruch nehmen. Natürlich ist es völlig anders, wenn man ein geselliges Gespräch führt oder im Ausland nach Informationen fragen möchte, aber im Fall von Gesundheitsproblemen, wo die genaue und präzise Beschreibung der Symptome für eine Diagnose unerlässlich ist, hätte ich mir viel lieber einen Dolmetscher oder eine Dolmetscherin gewünscht.

7. Nach der Untersuchung haben Sie gesagt, dass Sie sich auch im „realen Leben“ bei Problemen ohne jeden Vorbehalt zum Dr. [...] wenden würden, obwohl die meisten fachlichen Informationen mangelhaft und fehlerhaft von der App gedolmetscht wurden. Warum hatten Sie trotzdem den Eindruck, der Arzt wäre beruflich verlässlich und professionell?

Ich fand den Arzt sehr sympathisch, ich hatte den Eindruck, dass er trotz der schwierigen Kommunikation mein Grundproblem verstanden hat. Er führte die Untersuchung durch und konnte mir eine beruhigende Antwort bezüglich einer eventuell erblichen Krankheit geben. Er war fachlich sehr authentisch, hätte ich erstere Probleme, würde ich mich gerne an ihn wenden. Nach der Untersuchung habe ich sogar eine anatomische Aufklärung bzw. Tipps für eine gesunde Lebensführung bekommen, die ich aber in den meisten Fällen wegen der Fehlübersetzungen nicht wirklich verstanden habe.

8. In diesem Fall ging es über eine Vorsorgeuntersuchung, es gab keine ernsten Symptome. Hätten Sie echte und schmerzhaft Beschwerden gehabt, die weitere medizinische Behandlung oder sogar eine Operation benötigen, hätten Sie sich trotzdem getraut, sich der Untersuchung ohne einen Dolmetscher, nur mit Hilfe eines Tablets zu unterziehen?

Wie gesagt, würde ich mich gerne auch mit ernsteren, eventuell operationsbedürftigen Problemen an Dr. [...] wenden, ich würde mich aber auf keinen Fall mit Hilfe einer Applikation mit ihm unterhalten, sondern einen Dolmetscher beantragen.

Abstract

Abstract (Deutsch)

Das Ziel dieser Masterarbeit ist es, die Leistung der Dolmetsch-App *SayHi Translate* aufgabenorientiert, im Rahmen eines Arzt-Patient-Gesprächs zu evaluieren. Zu diesem Zweck wurde ein Experiment in einer Ordination mit einem österreichischen Gefäßchirurgen und einem ungarischen Patienten durchgeführt, die während einer Untersuchung verbal nur mit Hilfe der App kommunizieren konnten. Um einen Überblick über die historische Entwicklung und die Funktionsweise maschineller Dolmetsch-Systeme zu geben, werden im ersten Teil der Arbeit die Meilensteine aus der Geschichte des maschinellen Dolmetschens sowie die verschiedenen Komponenten solcher Systeme vorgestellt. Um die Problematik des DolmetscherInnenmangels im Gesundheitswesen zu verstehen, wird die Situation in Österreich in einem Kapitel beleuchtet. Nach der Experimentbeschreibung wird die Analyse durchgeführt, bei der vier verschiedene Aspekte (BenutzerInnenfreundlichkeit, Spracherkennungsfehler, Übersetzungsfehler und Usability – „Natürlichkeit der Interaktion“) berücksichtigt werden, um herauszufinden, inwieweit *SayHi Translate* dafür geeignet ist, die Kommunikation im Rahmen des Arzt-Patient-Gesprächs im Sprachenpaar Deutsch-Ungarisch zu ermöglichen. Anhand der Ergebnisse der Evaluierung lässt sich sagen, dass – nach derzeitigem Stand der Forschung – *SayHi Translate* nicht für die einwandfreie Kommunikation im medizinischen Bereich geeignet ist und professionelle DolmetscherInnen im Gesundheitswesen zurzeit nicht ersetzen kann.

Abstract (English)

The aim of this master's thesis is to conduct a task-oriented evaluation of the performance of the interpreting app *SayHi Translate* in the context of a conversation between a doctor and a patient. For this purpose, an experiment was carried out in a doctor's office with an Austrian vascular surgeon and a Hungarian patient who could communicate verbally only with the app during the examination. In order to provide an overview of the historical development and the functioning of machine interpreting systems, the first part of the thesis presents milestones in the history of machine interpreting and the various components of such systems. To understand the problem of the lack of interpreters in the healthcare system, the situation in Austria is examined in a chapter of the thesis. After the description of the experiment, the analysis is carried out. Four different aspects (user friendliness, speech recognition errors, translation errors and usability – "naturalness of the interaction") are taken into account to find out to what extent *SayHi Translate* can enable the communication between doctor and patient in the language pair German-Hungarian. Based on the results of the evaluation, it can be said that – according to the current state of research – *SayHi Translate* is not suitable for proper communication in the medical field and currently cannot replace professional interpreters in the healthcare sector.