



MASTERARBEIT | MASTER'S THESIS

Titel | Title

Maschinelles Dolmetschen mit Google Übersetzer
Möglichkeiten und Grenzen

verfasst von | submitted by

Pia Sommer BA BA BA BA

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

Wien | Vienna, 2026

Studienkennzahl lt. Studienblatt | Degree
programme code as it appears on the
student record sheet:

UA 070 331 345

Studienrichtung lt. Studienblatt | Degree
programme as it appears on the student
record sheet:

Masterstudium Translation Deutsch Französisch

Betreut von | Supervisor:

Univ.-Prof. Mag. Dr. Franz Pöchhacker

Inhalt

1 Einleitung	4
2 Begriffliche Grundlagen.....	6
2.1 Dolmetschen.....	6
2.2 Qualität und Erwartungen von Nutzer:innen	8
3 Maschinelles Dolmetschen.....	19
3.1 Funktionsweise.....	20
3.1.1 Spracherkennung	20
3.1.2 Maschinelles Übersetzen.....	22
3.1.3 Sprachausgabe	26
3.2 Forschungsstand	27
4 Methode.....	31
4.1 Datenerhebung	31
4.2 Datenanalyse	35
5 Ergebnisse	37
5.1 Spracheingabe	37
5.1.1 Rechtschreibung und Zeichensetzung.....	38
5.1.2 Verzögerungslaute und Selbstkorrekturen	41
5.1.3 Homophonie	42
5.1.4 Zusammenfassung	45
5.2 Maschinelle Übersetzung	46
5.2.1 Situatives Wissen und innertextliche Verweise	46
5.2.2 Lexikalische und syntaktische Ambiguitäten.....	47
5.2.3 Einfügungen und Auslassungen	48
5.2.4 Umformulierungen	49

5.2.5 Syntax.....	50
5.2.6 Verbesserungen	52
5.2.7 Zusammenfassung.....	53
5.3 Sprachausgabe	54
5.3.1 Ambiguitäten, Abkürzungen und Fremdwörter	54
5.3.2 Stimme	55
5.3.3 Zusammenfassung.....	56
5.4 Bewertung der maschinellen Dolmetschung.....	57
5.4.1 Übertragung des Inhalts	57
5.4.1.1 Treue zum Ausgangstext.....	57
5.4.1.2 Kohärenz	60
5.4.1.3 Pragmatik und Terminologie.....	61
5.4.2 Realisierung der Dolmetschung	62
5.4.2.1 Verzögerungslaute, Füllwörter und Sprechtempo.....	62
5.4.2.3 Grammatik.....	64
5.4.2.4 Syntax.....	65
5.4.2.5 Lexik und Stil	67
5.4.2.6 Länge	68
5.4.3 Quantitative Analyse	69
5.4.4 Gesamtbewertung.....	71
5.5 Bewertung der Humandolmetscher:innen	73
5.5.1 Übertragung des Inhalts	73
5.5.1.1 Auslassungen.....	73
5.5.1.2 Sinnverfälschungen	75
5.5.1.3 Umformulierungen	78

5.5.1.4 Kürzungen	80
5.5.1.5 Kohärenz, Kohäsion, Pragmatik und Terminologie.....	80
5.5.2 Realisierung der Dolmetschung	82
5.5.2.1 Stimme, Sprechtempo und Rhythmus.....	82
5.5.2.2 Aussprache, Intonation und Prosodie.....	83
5.5.2.3 Grammatik.....	84
5.5.2.4 Syntax.....	85
5.5.2.5 Lexik und Stil	86
5.5.3 Quantitative Analyse	87
5.5.4 Gesamtbewertung.....	93
6 Diskussion und Schlussfolgerungen.....	95
Bibliographie.....	105
Anhang	109
T 1 – Transkript 1 – Selbst erstelltes Transkript des Ausgangstextes.....	109
T 2 – Transkript 2 – Transkript des Ausgangstextes von Google Übersetzer.....	111
T3 – Transkript 3 –Übersetzung des Transkriptes von Google Übersetzer	113
T4 – Transkript 4 – Übersetzung des Transkriptes von Google Übersetzer inklusive der Pausenmarkierungen	115
T5 – Transkript 5 – Selbst erstelltes Transkript der Humandolmetschung 1	117
T6 – Transkript 6 – Selbst erstelltes Transkript der Humandolmetschung 2.....	119
T7 – Transkript 7 – Selbst erstelltes Transkript der Humandolmetschung 3.....	121
T8 – Transkript 8 – Selbst erstelltes Transkript der Humandolmetschung 4.....	123

1 Einleitung

Maschinelles Dolmetschen ist vor allem in den letzten Jahren immer mehr in den Fokus der translationswissenschaftlichen Forschung gerückt. Dies liegt nicht zuletzt daran, dass stets neue Software entwickelt wird und sich somit auch die Einsatzbereiche solcher Software immer mehr ausweiten und gleichzeitig die Qualität der Dolmetschleistungen solcher Programme verbessert werden soll (vgl. Fünfer 2014: 19ff.). Auch im laiensprachlichen Diskurs über Translation wird immer öfter die Frage aufgeworfen, ob es denn menschliche Translator:innen überhaupt noch braucht oder ob nicht die bereits bestehenden oder auch zukünftig noch erscheinenden Programme bereits eine so gute Leistung bieten können, dass Dolmetscher:innen dadurch ersetzt werden könnten. Aus diesem Grund werden solche Systeme auch von Translator:innen teilweise als existenzielle Bedrohung angesehen, da befürchtet wird, früher oder später gänzlich von Maschinen ersetzt zu werden (vgl. Fünfer 2014: 9).

Allerdings handelt es sich bei einer maschinellen Dolmetschung um einen äußerst komplexen Prozess, der aus drei einzelnen Schritten – der Spracherkennung, der maschinellen Übersetzung und der Sprachausgabe – besteht (vgl. Härtel 2016: 68), wodurch sich Fehlerquellen auf drei verschiedenen Ebenen ergeben können. Dazu zählen beispielsweise das fehlerhafte Erkennen von ähnlich klingenden Wörtern, das fehlende Kontextwissen für eine korrekte und adäquate Übertragung des Inhaltes in die Zielkultur oder Fehler bei der Aussprache von Homographen, Eigennamen oder Fremdwörtern (vgl. Härtel 2016: 79, 90, 97).

Andererseits könnten Programme zur maschinellen Dolmetschung in Zukunft auch eine wichtige Chance darstellen, sofern diese adäquate Leistungen erbringen. So bieten solche Programme etwa deutliche finanzielle Vorteile für Kund:innen, sind stets verfügbar und können somit auch spontan ohne Vorbereitung oder großen organisatorischen Aufwand eingesetzt werden. Demnach bieten sie sich gerade für einfache Alltagskommunikationssituationen, wie etwa in Urlaubssituationen, an, in denen normalerweise keine Dolmetscher:innen zur Verfügung stehen.

Aber können Programme zur maschinellen Dolmetschung wirklich mit den Leistungen von professionellen Dolmetscher:innen mithalten beziehungsweise in welcher Hinsicht und welchen Teilaspekten einer Dolmetschung sind diese Leistungen vielleicht sogar miteinander vergleichbar und wo ergeben sich auffallende Unterschiede hinsichtlich der Qualität und der Erfüllung der Erwartungen der Nutzer:innen?

Ausgehend von dieser Problematik ergibt sich folgende Forschungsfrage, die im Rahmen der vorliegenden Masterarbeit beantwortet werden soll: *Welche Möglichkeiten bietet die Software „Google Übersetzer“ und welche Grenzen im Vergleich zu menschlichen Dolmetscher:innen können aufgezeigt werden?*

Um diese Frage beantworten zu können, soll zunächst eine theoretische Einführung und Begriffserklärung der zentralen Konzepte erfolgen. In diesem Rahmen soll das Konzept des Dolmetschens behandelt werden, ehe auf Qualität beim Dolmetschen und die damit verbundenen Erwartungen seitens der Nutzer:innen eingegangen wird. Abschließend soll in diesem Kapitel noch erläutert werden, inwiefern Dolmetschleistungen bewertet werden können, damit dies für die spätere Analyse herangezogen werden kann. Im darauffolgenden Kapitel soll die Funktionsweise von maschinelltem Dolmetschen anhand der drei Zwischenschritte – Spracherkennung, maschineller Übersetzung und Sprachausgabe – beschrieben werden. Anschließend soll auf den aktuellen Forschungsstand rund um maschinelles Dolmetschen eingegangen werden, um einen Ausgangspunkt für die folgende Analyse erarbeiten zu können, ehe nach der Methodenbeschreibung die eigentliche Analyse der Dolmetschung von Google Übersetzer erfolgt. Dabei sollen die einzelnen drei Schritte der maschinellen Dolmetschung jeweils einzeln analysiert werden, um herausfinden zu können, bei welchem der Schritte sich jeweils Herausforderungen unterschiedlicher Art ergeben können. Danach soll auf die Bewertung der Leistung der Humandolmetscher:innen eingegangen werden, die nach dem gleichen Schema erfolgen soll, um abschließend den Vergleich zwischen Mensch und Maschine ziehen zu können, um zu eruieren, welche Möglichkeiten und Grenzen sich bei einer maschinellen Dolmetschung durch Google Übersetzer ergeben.

2 Begriffliche Grundlagen

In diesem Kapitel sollen zunächst grundlegende Begriffe definiert und erklärt werden, die für die weitere Masterarbeit von zentraler Bedeutung sind. Es handelt sich dabei einerseits um das Konzept des Dolmetschens an sich und andererseits um Qualität im Sinne der Translationswissenschaft und um die Bewertung von Dolmetschleistungen. Dies wiederum ist für die Durchführung der späteren Analyse und des Vergleichs sowie der Beurteilung der Dolmetschung der Software ‚Google Übersetzer‘ sowie jene der Humandolmetscher:innen wichtig, welche im weiteren Verlauf der Masterarbeit erfolgen wird.

2.1 Dolmetschen

Dolmetschen beschreibt grundsätzlich eine translatorische Handlung, die sich durch ihre besonderen Eigenschaften von anderen translatorischen Handlungen, wie etwa Übersetzen, abgrenzt. Einfach gesagt kann Dolmetschen als eine Art mündliches Übersetzen gesehen werden. Genauer betrachtet bedeutet Dolmetschen, dass ein Text übertragen wird, der nur einmal für die Dolmetscher:innen zur Verfügung steht und der folglich nicht wiederholt rezipiert werden kann. Demnach geschieht Dolmetschen stets unter einem gewissen Zeitdruck und in einem Live-Kontext (vgl. Pöchhacker 2010: 153f.).

Hinzu kommt, dass der Input für die Dolmetscher:innen auch multisemiotisch sein kann. Das bedeutet, dass neben den verbalen inhaltlichen Informationen auch nonverbale und paraverbale Eindrücke eine Rolle spielen können. Dazu zählen beispielsweise die Körpersprache der Sprecher:innen, ihr Verhalten im Raum oder die Prosodie (vgl. Pöchhacker 2010: 154). An dieser Stelle sei erwähnt, dass diese multisemiotischen Eindrücke von der Spracherkennung der Software zum maschinellen Dolmetschen nicht erkannt werden können. Auf die möglichen Folgen, die sich daraus ergeben können, wird in Kapitel 3.1 genauer eingegangen werden.

Die Tätigkeit des Dolmetschens kann dabei noch weiter unterteilt werden, wobei man zwischen dem Dolmetschmodus und dem Dolmetschsetting unterscheidet. In Bezug auf den Dolmetschmodus wird zwischen Konsekutiv- und Simultandolmetschen unterschieden, während bezüglich des Settings grundsätzlich zwischen Konferenz- und Dialogdolmetschen unterschieden wird (vgl. Dam 2010: 75; Pöchhacker 2010: 154f.). Da für den weiteren Verlauf dieser

Arbeit lediglich das Konsekutiv- sowie das Konferenzdolmetschen relevant sind, wird im Folgenden nur auf diese zwei Begriffe genauer eingegangen.

Konsekutivdolmetschen bedeutet, dass Dolmetscher:innen zunächst einen Ausgangstext anhören und den Zieltext erst im Anschluss daran produzieren. Folglich besteht der Prozess des Konsekutivdolmetschens aus zwei einzelnen aneinandergereihten Phasen, nämlich der Rezeption – also dem Zuhören und Verstehen – und der Produktion. Die zu übertragenden Passagen können dabei in ihrer Länge variieren und reichen von wenigen Wörtern bis hin zu ganzen Reden, die bis zu zehn Minuten umfassen können (vgl. Ahrens 2016: 84f.; Dam 2010: 75).

Aufgrund des natürlichen Ablaufs dieses Dolmetschprozesses, für den keinerlei technische Hilfsmittel benötigt werden, gilt Konsekutivdolmetschen als älteste und ursprüngliche Form der Translation. Konsekutivdolmetschen kann dabei sowohl monologisch, also unilateral, als auch dialogisch, also bilateral, zur Anwendung kommen (vgl. Ahrens 2016: 84, 91; Dam 2010: 75f.). Welche Form des Konsekutivdolmetschens verwendet wird, hängt dabei stark vom jeweiligen Setting ab, worauf im weiteren Verlauf dieses Kapitels noch genauer eingegangen wird.

Durch diese zwei Phasen des Dolmetschprozesses liegt der Vorteil des Konsekutivdolmetschens darin, dass den Dolmetscher:innen die Möglichkeit gegeben wird, länger über gewisse Formulierungen nachzudenken. Das wiederum kann zu einer besseren sprecherischen Leistung der Dolmetscher:innen führen und somit eine sprachlich einwandfreie Wiedergabe begünstigen (vgl. Ahrens 2016: 86; Setton 2010: 69). Außerdem kommt es beim Konsekutivdolmetschen aufgrund der fehlenden technischen Ausrüstung seltener zu technischen Störungen und auch der Einsatz von Konsekutivdolmetscher:innen ist deswegen oft billiger, da für keine technische Ausrüstung bezahlt werden muss (vgl. Ahrens 2016: 90).

Im Gegensatz dazu liegt der Nachteil beim Konsekutivdolmetschen darin, dass mehr Zeit dafür aufgewendet werden muss, weshalb der Einsatz normalerweise nur bei kürzeren Ausgangstexten, die nicht gleichzeitig in mehrere Sprachen gedolmetscht werden sollen, praktikabel ist. Dies ist auch der Grund, weshalb das Konsekutivdolmetschen im Rahmen von Konferenzen immer mehr vom Simultandolmetschen abgelöst wurde und heutzutage in einem solchen Setting nur mehr selten zum Einsatz kommt¹ (vgl. Ahrens 2016: 92; Dam 2010: 76; Setton 2010: 69).

¹ Aufgrund der begrenzten Möglichkeiten der hier untersuchten Software ‚Google Übersetzer‘, welche aktuell lediglich eine Konsekutivdolmetschung und noch keine Simultandolmetschung eines längeren Ausgangstextes erlaubt, wird im Rahmen dieser Arbeit trotzdem das Konsekutivdolmetschen als Dolmetschmodus für die Human-dolmetscher:innen herangezogen, um einen besseren Vergleich zu ermöglichen.

Das bereits kurz angesprochene Konferenzdolmetschen bezeichnet folglich also nur das Setting, weshalb grundsätzlich beide Dolmetschmodi hierbei verwendet werden können, wobei hier meistens nur eine unilaterale Dolmetschung zum Einsatz kommt, da in einem solchen Kontext häufig Vorträge beziehungsweise Reden und eher selten dialogische Gespräche gedolmetscht werden (vgl. Kalina 2012: 134). Konferenzdolmetschen ist dabei ein sehr professionalisiertes Feld, was sich vermutlich auch durch die formelle Situation und die damit einhergehende stark standardisierte Rolle der Dolmetscher:innen ergibt (vgl. Pöchhacker 2010: 156; Setton 2010: 66f.). Das bedeutet, dass Konferenzdolmetscher:innen meist professionell und universitär ausgebildet sind und oft auch internationalen Normen, Konventionen und Standards unterliegen. Diese werden dabei etwa auch von Berufsverbänden vorgegeben. Der weltweit größte Berufsverband ist dabei die AIIC (Association Internationale des Interprètes de Conférence), der internationale Konferenzdolmetscherverband, welcher technische, ethische und qualitative Standards für Konferenzdolmetscher:innen vorgibt (vgl. Setton 2010: 67f., 71). Diese Vorgaben dienen wiederum der Qualitätssicherung der Dolmetschleistungen, worauf im folgenden Kapitel genauer eingegangen wird.

2.2 Qualität und Erwartungen von Nutzer:innen

In diesem Kapitel soll der Begriff Qualität in einem translationswissenschaftlichen Kontext definiert und davon ausgehend auch Qualitätskriterien festgelegt werden, welche wiederum die Grundlage für das Bewerten von Dolmetschleistungen bilden, auf das im nächsten Kapitel genauer eingegangen wird. Außerdem sollen die Erwartungen, die seitens der Nutzer:innen an eine Dolmetschleistung gestellt werden, behandelt werden.

Die Bewertung der Dolmetschqualität gilt als eine der Schlüsselfragen innerhalb der Dolmetschwissenschaft, da diese einen wichtigen Beitrag zur Qualitätssicherung von Dolmetschungen leistet (vgl. Kutz 2005: 14). Dies ist besonders wichtig, da Dolmetscher:innen als Dienstleister:innen agieren, die eine Kommunikation zwischen mehreren Parteien ermöglichen, die anderenfalls gar nicht oder nur sehr eingeschränkt möglich wäre (vgl. Pöchhacker 2015: 430). Das bedeutet allerdings auch, dass die meisten Nutzer:innen einer Dolmetschleistung keinerlei Kontrolle über deren Qualität haben, wodurch Dolmetschleistungen oft als notwendiges Übel angesehen werden. Dies liegt daran, dass Nutzer:innen einer Dolmetschung deren Korrektheit normalerweise nicht nachprüfen können. Demnach treten Nutzer:innen

Dolmetscher:innen auch oft mit Skepsis gegenüber, da Dolmetschfehler zu großen Missverständnissen und peinlichen Situationen führen können (vgl. Kalina 2004: 2).

Als Nutzer:innen werden hierbei in erster Linie jene Personen verstanden, die an einer Kommunikationssituation mit einer Dolmetschung direkt beteiligt sind. Dies betrifft folglich einerseits die Ausgangstextproduzent:innen und andererseits auch die Zieldextrezipient:innen. Im weiteren Sinne könnten zu Nutzer:innen auch eventuelle Auftraggeber:innen, Agenturen oder ähnliche Akteur:innen zählen, die im Laufe des Prozesses von der Erteilung eines Dolmetschauftrages bis hin zur Durchführung mitarbeiten (vgl. Pöchhacker 2015: 430). Im Zuge dieser Arbeit liegt der Fokus allerdings auf der direkt an der Kommunikation beteiligten Nutzer:innengruppe.

Da die Qualitätssicherung also eine wichtige Rolle in der Dolmetschwissenschaft spielt, ist es von zentraler Bedeutung, das Konzept ‚Qualität‘ zu definieren, um es schließlich auch auf die Bewertung von Dolmetschleistungen anwenden zu können. Dabei ist jedoch zu beachten, dass Qualität grundsätzlich ein eher schwer zu erfassendes, komplexes Konzept darstellt, welches nicht einfach allgemeingültig definiert werden kann, sondern aus verschiedenen ineinandergreifenden Faktoren besteht, die sich je nach Gegebenheiten voneinander unterscheiden können (vgl. Christensen 2001: 1; Shlesinger et al. 1997: 123).

Vereinfacht gesagt, kann Qualität dabei prinzipiell als das grundlegende Erreichen von Zielen angesehen werden (vgl. Grbić 2008: 236; Shlesinger et al. 1997: 127). Damit geht einher, dass zur Verwirklichung dieser Ziele bestimmte Erwartungen erfüllt werden müssen. Dabei stellt sich folglich die Frage, welche Ziele und Erwartungen es in Bezug auf Dolmetschleistungen gibt, die es zu erfüllen gilt, worauf im weiteren Verlauf dieses Kapitels noch genauer eingegangen werden wird.

Weiters gilt somit, dass der Grad der Qualität daran gemessen werden kann, inwieweit geforderte beziehungsweise erwartete Standards oder Normen erfüllt werden. Dies kann in Bezug auf Dolmetschen dabei grundsätzlich aus zwei Perspektiven – aus der linguistischen und der pragmatischen Perspektive – betrachtet werden. Als Pragmatik wird hierbei das zielgerichtete Verwenden von Sprache verstanden, wodurch ein spezieller Effekt erzielt werden kann. Der linguistische Sinn bezieht sich dabei auf den Inhalt sowie die Form, die Ausgangs- und Zieldtext aufweisen (vgl. Kopczyński 1994: 87f.).

Qualität ist somit stets relativ zu betrachten, da Qualität weder als absolut zu sehen ist noch allgemeingültig definiert werden kann, sondern immer abhängig vom Kontext und den jeweiligen Beurteiler:innen und ihren individuellen Erwartungen und Anforderungen ist (vgl. Gile 2001: 380). Folglich bedeutet dies, dass eine Vielzahl von voneinander abhängigen

Variablen in ihrem Zusammenspiel betrachtet werden müssen, um den Grad der Qualität feststellen zu können. So besteht das Konzept von Qualität darin, dass unterschiedlich gewichtete Qualitätsparameter in ihrer Gesamtheit als Qualitätsmaßstab fungieren. Demnach ist Qualität in der Dolmetschwissenschaft auch eng mit dem jeweiligen Skopos einer Dolmetschung verbunden, da zur Erfüllung des jeweiligen Skopos auch die einzelnen Qualitätskriterien beachtet werden müssen (vgl. Grbić 2008: 237ff.).

Wenn nun Erwartungen an Dolmetschleistungen eine wichtige Rolle zum Erreichen von Zielen und damit einhergehend zur Findung von Qualitätskriterien spielen, kann die Zufriedenheit der Nutzer:innen einer Dolmetschleistung somit als eine Art Qualitätsparameter herangezogen werden. Das bedeutet, dass das Verhältnis zwischen den Erwartungen der Nutzer:innen und der tatsächlich erbrachten Leistung betrachtet wird, wobei gilt, dass die Erwartungen zumindest eingehalten oder im besten Fall sogar übertroffen werden sollen, damit die Nutzer:innen zufrieden sind. Denn grundsätzlich gilt, dass die Dolmetschleistung als qualitativ hochwertig angesehen werden kann, wenn alle Nutzer:innen einer Dolmetschung mit der erbrachten Leistung zufrieden sind. Anders gesagt bedeutet Qualität somit in der Dolmetschwissenschaft, dass die Dienstleistung, die Dolmetscher:innen erbringen, den Anforderungen und den Erwartungen der Kund:innen beziehungsweise Nutzer:innen entsprechen sollte (vgl. Christiansen 2011: 2; Gouadec 2010: 271; Kurz 2001: 404f.).

Demnach sind in Hinblick auf die Qualitätskriterien und die Erfüllung von Zielen einer Dolmetschleistung zunächst die Erwartungen der Nutzer:innen einer Dolmetschleistung wichtig, da es die grundsätzliche Aufgabe von Dolmetscher:innen ist, eine erfolgreiche und effektive Kommunikation zu ermöglichen, die in der jeweiligen Situation adäquat ist. Dabei ist festzuhalten, dass es Rezipient:innen einer Dolmetschung aufgrund der mangelnden Sprachkompetenz in der Ausgangssprache meistens nicht möglich ist, einen Vergleich mit dem Original zu ziehen. Demnach können diese etwa selten wirklich beurteilen, ob eine Dolmetschung vollständig ist oder ob Sinnübereinstimmung mit dem Original besteht. Andererseits können Rezipient:innen den Zieltext basierend auf der Stimme, dem Akzent und der Sprachkompetenz in der Zielsprache beurteilen (vgl. Bühler 1986: 233).

Nichtsdestotrotz empfinden Nutzer:innen einer Dolmetschung unter anderem Sinnübereinstimmung des Zieltextes mit dem Ausgangstext, Kohärenz des Zieltextes und eine inhaltlich vollständige Dolmetschung als wichtige Kriterien, die es zu erfüllen gilt, um eine als qualitativ hochwertig empfunden Dolmetschleistung zu liefern (vgl. Kurz 2001: 406). Außerdem spielt die Korrektheit der verwendeten Terminologie eine wichtige Rolle. Hinzu kommen eine lebhaft und nicht monoton klingende Stimme und vollständige, flüssig und klar formulierte Sätze,

die ohne Verzögerungslaute und übermäßige Pausen auskommen, sowie Mikrofondisziplin, so dass möglichst keine Störgeräusche, wie etwa Husten oder Räuspern auftreten (vgl. Moser 1995: 14, 26-29). Auf der anderen Seite werden Grammatikfehler oder ein anderssprachlicher Akzent der Dolmetscher:innen von Nutzer:innen als nicht so störend empfunden (vgl. Kurz 2001: 406).

Dabei stellt sich jedoch die Frage, ob Nutzer:innen selbst wirklich wissen, was sie von einer qualitativ hochwertigen Dolmetschung erwarten können. Aus diesem Grund ist es wichtig, neben der Perspektive der Erwartungen der Nutzer:innen auch andere Aspekte in die Qualitätsbewertung miteinzubeziehen. Dazu gehören etwa weitere Erkenntnisse aus der translati-onswissenschaftlichen Forschung, wie beispielsweise eine korrekte Sinnübertragung, die bestenfalls stattfinden soll (vgl. Shlesinger et al. 127: 126). Diese Faktoren werden dabei in weiterer Folge und auch im folgenden Kapitel noch genauer behandelt und erläutert.

Außerdem ist es wichtig zu beachten, dass es seitens der Nutzer:innen keine einheitlichen und allgemeingültigen Erwartungen an Dolmetschleistungen gibt, sondern dass sich diese teilweise stark voneinander unterscheiden können. Dabei spielen nicht nur die persönliche Meinung von befragten Personen eine Rolle, sondern auch weitere Faktoren, wie beispielsweise das Setting, in dem eine Dolmetschung stattfindet. Hinzu kommen zusätzliche Faktoren wie der Dolmetschmodus oder auch das Thema sowie der Grad der Fachlichkeit des Ausgangstextes. Bezüglich der Erwartungen der Nutzer:innen von Dolmetschungen lässt sich folglich feststellen, dass diese je nach Setting oder Kontext unterschiedlich aussehen können, wodurch sich wiederum unterschiedliche Qualitätskriterien ergeben. Dabei ist außerdem zu beachten, dass nicht ein Merkmal allein als Qualitätskriterium herangezogen werden kann, sondern stets mehrere Kriterien ineinandergreifen und somit als Gesamteindruck betrachtet werden sollten (vgl. Kalina 2004: 3, 5, 2012: 135).

Das bedeutet weiters, dass es sich bei Qualität nicht um einen absoluten Wert handeln kann, sondern Qualität stets als relativer und stark kontextabhängiger Wert gesehen werden muss (vgl. Kopczyński 2011: 88). So sind die jeweiligen Qualitätskriterien und -maßstäbe stets abhängig von den betreffenden Akteur:innen, die jeweils gewisse Qualitätsstandards erwarten (vgl. Grbić 2008: 240). Diese Qualitätsstandards können sich dabei auf die unterschiedlichsten Aspekte einer Dolmetschleistung beziehen, wozu etwa Semantik, Prosodie, Pragmatik, Lexik, Stimme oder auch Syntax gehören (vgl. Gile 2001: 380).

Trotz dieser unterschiedlichen Erwartungen und Ansprüche bezüglich Qualität besteht eine ungefähre Einigkeit, welche sich hierbei lediglich auf grundlegende Erwartungen an eine Dolmetschleistung bezieht. Dazu zählen etwa Faktoren wie beispielsweise

Sinnübereinstimmung, Kohärenz, korrekte Terminologie, Fehlerfreiheit, Flüssigkeit oder auch Treue zum Original. All diese Faktoren variieren dabei jedoch wiederum in ihrer jeweiligen Relevanz je nach den bereits genannten Faktoren wie etwa Setting oder Thema der Dolmetschung (vgl. Kalina 2012: 135; Pöchhacker 2001: 413, 2015: 431).

Als eines der wichtigsten Ziele wird die Treue zum Ausgangstext genannt. Dies wird ebenso in der Dolmetschwissenschaft als Hauptziel einer qualitativ hochwertigen Dolmetschung angesehen und ist besonders wichtig, da dadurch erst eine effektive Kommunikation ermöglicht wird, was wiederum das grundlegende Ziel von Dolmetschleistungen darstellt. Diese Treue zum Ausgangstext bedeutet grundsätzlich, dass der Zieltext, den die Rezipient:innen einer Dolmetschung hören, denselben Effekt auf diese haben sollte wie der Ausgangstext auf dessen direkte Rezipient:innen hat. Das bezieht sich dabei etwa auf den Inhalt, die Wirkung, den Stil oder auch den Informationsgehalt (vgl. Kurz 2001: 394f.). Dieser Qualitätsparameter ist dabei allerdings nicht einfach mit einer möglichst wörtlichen Ähnlichkeit zum Ausgangstext gleichzusetzen, da hier weitere Faktoren, wie die Unterschiede zwischen den verwendeten Sprachen und Kulturen, bedacht werden müssen (vgl. Gile 2009: 52). Wie eine Bewertung dieses Faktors genau aussieht, wird im Kapitel 2.3 im Kontext der konkreten Bewertung von Dolmetschleistungen genauer erläutert.

Neben der Treue zum Ausgangstext gibt es – wie bereits kurz angesprochen – einige weitere Faktoren, die als gemeinsame Erwartungen der meisten Nutzer:innengruppen an qualitativ hochwertige Dolmetschleistungen gestellt werden. Diese gehen dabei meist ebenfalls mit einer gewissen Äquivalenz zwischen Ausgangs- und Zieltext einher, weshalb weitere inhaltliche Aspekte wie terminologische Korrektheit und Exaktheit ebenso als wichtig für eine qualitativ hochwertige Dolmetschung erachtet werden. Die Wichtigkeit der terminologischen Korrektheit variiert dabei klarerweise je nach Grad der Fachlichkeit eines Settings. Das bedeutet, dass je fachlicher ein Ausgangstext ist und je mehr Fachterminologie darin enthalten ist, desto wichtiger wird auch die korrekte Verwendung der Terminologie sowohl von den Sender:innen als auch von den Rezipient:innen angesehen, da die Terminologie einen wichtigen Betrag zur korrekten Übertragung des Inhalts und zu einer besseren Verständlichkeit dessen beiträgt (vgl. Kopczyński 2011: 93). Außerdem spielen weitere sprachliche Faktoren wie beispielsweise korrekte Grammatik und die Kohärenz des Zieltextes (vgl. Pöchhacker 2004: 173f.) eine wichtige Rolle, da auch diese dazu beitragen, den Inhalt des Ausgangstextes für die Zuhörer:innen verständlich wiederzugeben.

Neben den bereits genannten inhaltlichen Aspekten einer Dolmetschleistung gelten prosodische Parameter, wie beispielsweise Intonation, Flüssigkeit, Akzent, oder auch die generelle

Wahrnehmung der Stimme der Dolmetscher:innen ebenso als wichtige Qualitätsparameter. Dabei ist jedoch zu beachten, dass diese oftmals als weniger relevant als inhaltliche Aspekte betrachtet werden, wobei Mängel an solchen rhetorischen Fähigkeiten wiederum von den meisten Nutzer:innen als störend empfunden werden (vgl. Collados Aís 2016: 676, 684; Moser 1995: 21). Dies lässt sich wohl vor allem damit erklären, dass unter anderem monotones Sprechen, eine unpassende Sprechgeschwindigkeit und weitere paraverbale Parameter zu Ermüdung oder Ablenkung bei den Rezipient:innen führen, wodurch in weiterer Folge wiederum das inhaltliche Verständnis des Zieltextes beeinträchtigt werden könnte (vgl. Collados Aís 2016: 684f.).

Diese Qualitätsfaktoren, die sich nicht direkt auf den inhaltlichen Aspekt einer Dolmetschung beziehen, lassen jedoch wiederum erkennen, welche wichtige Rolle der inhaltliche Aspekt bei der Festlegung von Qualitätskriterien spielt, da durch solche nonverbalen beziehungsweise paraverbalen Faktoren auch inhaltliche und verständnisbezogene Faktoren sowohl positiv als auch negativ beeinflusst werden können. Demnach ist es wichtig, all diese unterschiedlichen Faktoren gemeinsam und in ihrem Zusammenspiel zu betrachten, um eine Dolmetschung als qualitativ hochwertig bewerten zu können. Wie eine solche umfassende Bewertung konkret aussehen kann, wird im folgenden Kapitel genauer erläutert.

In Bezug auf all die bereits genannten Qualitätsfaktoren ist außerdem zu beachten, dass diese nicht alle in der Macht der Dolmetscher:innen selbst liegen, sondern es auch externe Faktoren gibt, die die Qualität einer Dolmetschleistung beeinflussen können. Dazu zählt etwa die konkrete Dolmetschsituation mit den damit einhergehenden Bedingungen. Das bedeutet, dass beispielsweise die Qualität der technischen Ausrüstung – Mikrofone, Kopfhörer oder ähnliches – das Redetempo, die inhaltliche Dichte und die grammatikalische Struktur des Ausgangstextes oder auch seitens der Auftraggeber:innen beziehungsweise der Veranstalter:innen zur Verfügung gestellte Zusatzmaterialien die Qualität einer Dolmetschleistung stark beeinflussen können (vgl. Kalina 2004: 6, 2012: 137).

In Bezug auf das Konzept der Qualität innerhalb der Dolmetschwissenschaft und die damit einhergehenden Qualitätsparameter beziehungsweise -kriterien lässt sich also zusammenfassend feststellen, dass es keine allgemeingültige Definition von Qualität gibt, sondern dass diese vielmehr von der jeweiligen Situation abhängt. Dabei ist es jedoch stets wichtig, die jeweiligen Erwartungen und Ziele zu erfüllen, um von einer qualitativ hochwertigen Dolmetschleistung sprechen zu können. Dabei spielen folglich erwartete Qualitätskriterien eine Rolle, die zusammengefasst bedeuten, dass der Kommunikationsprozess, der durch Dolmetscher:innen ermöglicht wird, erfolgreich ablaufen muss, damit eine Dolmetschleistung als qualitativ hochwertig angesehen werden kann (vgl. Pöchhacker 2001: 413).

2.3 Bewerten von Dolmetschleistungen

In diesem Kapitel soll erklärt werden, wie Dolmetschleistungen bewertet werden können. Dies erfolgt grundsätzlich aufbauend auf den Qualitätskriterien und den Erwartungen, die im vorangegangenen Kapitel erläutert wurden. Dabei sollen diese Kriterien einerseits nach inhaltlichen Gemeinsamkeiten und somit Bewertungsebenen und andererseits nach Relevanz für die später folgende Analyse geordnet werden. Dadurch soll ein Schema entstehen, anhand dessen in der späteren Masterarbeit die entsprechenden Dolmetschleistungen der Humandolmetscher:innen und der Software ‚Google Übersetzer‘ nicht nur miteinander verglichen, sondern auch bewertet werden können.

Laut Kutz (2005) lassen sich die Bewertungskriterien für das Konsekutivdolmetschen in drei Gruppen unterteilen. Dazu zählt der Eindruck beziehungsweise das Verhalten der Dolmetscher:innen, die Inhaltswiedergabe und die Form der Realisierung (vgl. Kutz 2005: 21f., 24). Eine weitere Unterteilung der Qualitätsparameter, die auch gut für die Bewertung von Dolmetschleistungen herangezogen werden kann, unterscheidet zwischen Intertextualität, Intratextualität und Instrumentalität (vgl. Shlesinger et al. 1997: 128). Ähnliche Punkte werden bei Han (2018) angeführt, wo als die drei Hauptkriterien zur Bewertung von Dolmetschleistungen der Inhalt, die Übermittlung und die Sprachqualität genannt werden. Diese können durch weitere Kriterien, wie Aussprache, Grammatik oder auch Informationsgehalt, ergänzt werden, um eine umfassende Bewertung zu gewährleisten (vgl. Han 2018: 64).

Diese genannten Schemata zur Bewertung beziehen sich inhaltlich auf ähnliche Qualitätsparameter, die lediglich auf unterschiedliche Weise geordnet werden. Deshalb werden im weiteren Verlauf dieses Kapitels die Ansätze miteinander kombiniert und mit weiteren Qualitätsparametern ergänzt, um einheitliche Kriterien für die spätere Analyse zu schaffen. Das Schema für die Einordnung der Kriterien wird sich dabei vom Aufbau an dem dreigeteilten Schema von Kutz (2005) orientieren. Dabei ist außerdem wichtig zu beachten, dass das Bewerten von Dolmetschleistungen stets auf möglichst objektive Art und Weise erfolgen muss (vgl. Colina 2011: 43), weshalb der Fokus darauf liegen sollte, dass alle geforderten Kriterien im gewollten Maß erfüllt werden.

Als wichtigste Grundvoraussetzung für eine gelungene Dolmetschung gilt, wie auch bereits im vorangegangenen Kapitel erwähnt, die Übertragung des Inhalts, also die korrekte Inhaltswiedergabe. Hierbei wird oft von einer gewissen Treue zum Ausgangstext gesprochen. Diese bezieht sich dabei allerdings nicht auf eine wörtliche Treue, sondern vielmehr darauf,

dass die Empfänger:innen der Dolmetschung die Nachricht des Ausgangstextes genauso gut wie direkte Rezipient:innen des Ausgangstextes verstehen (vgl. Gile 2009: 37). Demnach ist es wichtig, den Ausgangstext mit dem Zieltext zu vergleichen und in Bezug auf mögliche Auslassungen Fehler, Widersprüche oder ähnliches zu überprüfen, um zu erkennen, ob die gewünschte Treue erreicht wurde. Dabei müssen selbstverständlich grammatikalische und lexikalische Unterschiede, die zwischen den untersuchten Sprachen bestehen, bedacht werden (vgl. Gile 2009: 52).

Folglich bedeutet dies, dass das Hinzufügen oder Auslassen von einzelnen Wörtern oder Phrasen allein noch nicht zwangsläufig etwas an dem Grad der Treue zum Ausgangstext ändert, da eine Nachricht prinzipiell unterschiedlich verbalisiert werden kann. Das bedeutet, dass unterschiedliche Sätze auch innerhalb einer Sprache dieselbe Nachricht mit demselben Informationsgehalt ausdrücken können. Dies betrifft dann klarerweise auch Sätze in verschiedenen Sprachen (vgl. Gile 2009: 53, 56).

Gile (2009) spricht in diesem Zusammenhang von „information gains“ und „information loss“ (Gile 2009: 56). Damit ist jeweils das Hinzufügen beziehungsweise das Weglassen von Informationen gemeint. Auch hierbei ist zu beachten, dass es sich etwa bei Umformulierungen oder auch zusätzlichen Erklärungen seitens der Dolmetscher:innen nicht zwangsläufig um solche ‚information gains‘ handelt, wenn diese das Verständnis der Zieltextrezipient:innen verbessern. Gleiches gilt ebenso für das Weglassen von redundanten oder für das Zielpublikum bereits bekannten Informationen. Auch dies bedeutet nicht zwingend das Vorhandensein von ‚information loss‘ (vgl. Gile 2009: 56f., 62). Ebenso können Dolmetscher:innen eindeutige inhaltliche Fehler, die etwa auf Versprechern seitens der Ausgangstextproduzent:innen beruhen, ausbessern, ohne dass dabei von einem Verstoß gegen der Treue zum Ausgangstext gesprochen werden kann (vgl. Gile 2009: 63).

Wenn nun im Zuge der Bewertung des inhaltlichen Aspektes von Treue zum Ausgangstext gesprochen wird, ist wichtig zu beachten, dass es stets darum geht, den gleichen Zweck beziehungsweise den gleichen Eindruck bei den Rezipient:innen hervorzurufen, der auch von den Ausgangstextproduzent:innen angestrebt wurde (vgl. Colina 2011: 44f.).

In Bezug auf diese inhaltlichen Aspekte spielt aber nicht nur die reine Treue zum Ausgangstext eine wichtige Rolle für die Bewertung, sondern auch weitere Faktoren beziehungsweise besonders wichtige Punkte, die auch bei der Übertragung des Inhalts zu beachten sind. Dazu zählt die Kohärenz des Dolmetschproduktes, also die logische Stimmigkeit, sowie die Verwendung von angemessenen Verknüpfungen. Außerdem ist es wichtig, dass sogenannte pragmatische Sensibilitäten exakt wiedergegeben werden und bei möglichen

Verständnisschwierigkeiten für das Zielpublikum entsprechend erklärt werden. Dazu zählen Eigennamen, Positionen, Titel oder auch Berufsbezeichnungen (vgl. Kutz 2005: 22ff.). Weiters ist, wie bereits im vorangegangenen Kapitel erwähnt, die korrekte Verwendung der Terminologie sehr bedeutend, da Fehler, die in diesem Zusammenhang auftreten, das inhaltliche Verständnis des Zieltextes stark beeinträchtigen können (vgl. Bühler 1986: 232; Moser 1995: 16).

Neben der Bewertung der inhaltlichen Ebene gilt die Bewertung der Realisierung der Dolmetschung ebenfalls als relevant. Das bedeutet, dass die Dolmetschung auch auf funktionalstilistischer Ebene angemessen produziert werden sollte. Dazu zählen Faktoren, wie die Stimme, das Sprechtempo, die Rhythmik und im weiteren Sinne auch die Sprachkompetenz, die Aussprache, sowie die korrekte Verwendung von Grammatik, Syntax, Lexik und Stil. Für das Konsekutivdolmetschen relevant ist auch der zeitliche Faktor. Eine Konsekutivdolmetschung sollte in Bezug auf die Länge demnach in etwa 70% des Originals in Anspruch nehmen, um einerseits zwar keine groben inhaltlichen Einbußen hinnehmen zu müssen und andererseits aber auch die Aufmerksamkeitsspanne der Rezipient:innen nicht zu übersteigen und nicht zu viel Zeit in Anspruch zu nehmen, da diese gerade in einem Konferenzsetting oftmals knapp bemessen ist (vgl. Kutz 2005: 25ff.).

Besonderes Augenmerk soll hierbei ebenso auf die Intonation gelegt werden, da eine angenehme Stimme mit gut eingesetzter Intonation dabei helfen kann, den Inhalt besser verständlich zu vermitteln und Verständnisproblemen vorbeugen zu können (vgl. Collados Aís 2016: 677). Weiters ist es wichtig, dass eine angenehme Sprechgeschwindigkeit eingehalten wird, da zu schnelles Sprechen das Verständnis beeinträchtigen kann, während zu langsames Sprechen zu schnellerer Ermüdung beziehungsweise Ablenkung bei den Rezipient:innen führen kann. Selbiges gilt für zu monotones Sprechen, da durch die entstehende Langeweile und die erhöhte Anstrengung beim Zuhören ebenso das Verständnis des Inhaltes erschwert werden kann (vgl. Collados Aís 2016: 684f.).

Als drittes übergeordnetes Bewertungskriterium gelten der generelle Eindruck und das Verhalten der Dolmetscher:innen. Das betrifft dabei beispielsweise die Körperhaltung, die Körpersprache, den Blickkontakt, die Gestik oder auch die Mimik. Allgemein ist hierbei wichtig, dass die Dolmetscher:innen eine angemessene Professionalität während ihres Berufseinsatzes ausstrahlen (vgl. Kutz 2005: 21f.). Diese Kriterien werden allerdings im Laufe dieser Arbeit nicht näher erläutert, da diese für die spätere Analyse nur von marginaler Bedeutung sind. Dies liegt daran, dass sich diese Analyse und damit der Vergleich zwischen der Dolmetschleistung der Software und jener der Humandolmetscher:innen auf Audioaufnahmen und daraus angefertigte Transkripte stützen wird. So soll der Fokus auf der sprachlichen und prosodischen

Bewertungsebene liegen, um die Dolmetschleistungen untereinander besser vergleichbar zu machen, worauf in Kapitel 4 noch genauer eingegangen wird.

Neben den bereits genannten gewünschten Anforderungen an die Dolmetschleistungen soll nun noch auf eventuell vorkommende Fehler eingegangen werden, um diese ebenfalls einordnen und genauer erläutern zu können. Dabei lässt sich zunächst feststellen, dass nicht immer eindeutig erkennbar ist, wann es sich um einen Fehler handelt und welche möglichen Abweichungen noch als akzeptabel gelten können. Hierbei ist – wie bei den meisten Qualitätskriterien auch – stets die konkrete Situation und der jeweilige Kontext in die Bewertung einzubeziehen (vgl. Gile 2001: 387). Demnach gilt es auch hier, den Fokus nicht etwa auf einzelne Wörter oder ähnliches zu richten, sondern die Dolmetschleistung stets im Kontext und als Gesamtleistung zu betrachten.

Nichtsdestotrotz können im Kontext gesehen einzelne Fehler erkannt werden, die sich folglich klarerweise negativ auf die Gesamtbewertung auswirken. Dabei können sich diese Fehler wiederum auf alle bereits genannten Bewertungsebenen beziehen. Kopczyński (1980) unterscheidet dabei vier Arten von Fehlern, wozu Translationsfehler, Sprachkompetenzfehler, Fehler der sprachlichen Leistung und Fehler der kommunikativen Angemessenheit gehören. Translationsfehler beziehen sich dabei auf inhaltliche Fehler und somit auf Mängel bezüglich der bereits angesprochenen Treue zum Ausgangstext. Demnach sind diese Fehler besonders schwerwiegend, da hier die Nachricht des Ausgangstextes verfälscht wiedergegeben wird. Als Sprachkompetenzfehler gelten jene Fehler, die zum Sprachvermögen der Dolmetscher:innen zählen, wie etwa grammatikalische oder syntaktische Fehler. Diese Fehler wiederum können – im Gegensatz zu den Translationsfehlern – von den Rezipient:innen besonders leicht erkannt werden. Deshalb können diese Fehler als störend empfunden werden, während sie nicht zwangsläufig das Inhaltsverständnis beeinträchtigen. Die Fehler der sprachlichen Leistung beziehen sich auf Performanzprobleme, wozu etwa Stottern, Pausen oder auch die häufige Verwendung von Füllwörtern zählen. Fehler der kommunikativen Angemessenheit schließlich werden als Kombination aus mehreren Fehlern beziehungsweise Fehlertypen angesehen, die zu größeren Verständnisproblemen führen und somit eine reibungslose Kommunikation verhindern können (vgl. Kopczynski 1980: 291-296).

Zusammenfassend lässt sich zu der Bewertung von Dolmetschleistungen also feststellen, dass der Ausgangstext möglichst genau und vollständig wiedergegeben werden sollte, so dass keine wichtigen Informationen verloren gehen und vor allem die Botschaft dieses Textes nicht verzerrt wird. Dies ist die Grundvoraussetzung und das wichtigste Kriterium für eine gelungene und qualitativ hochwertige Dolmetschleistung. Weiters sollte der Zieltext sowohl

sprachlich als auch prosodisch angemessen von den Dolmetscher:innen produziert werden, um das Verständnis der Rezipient:innen zu erleichtern und das Zuhören möglichst angenehm zu gestalten. Insgesamt sollte durch eine solche Dolmetschleistung eine effektive und möglichst problemlose Kommunikation ermöglicht werden (vgl. Fünfer 2013: 30). Dabei gilt es ebenso etwaige Fehler zu vermeiden, welche wiederum eine negative Auswirkung auf die Gesamtbewertung einer Dolmetschleistung haben können. Folglich ist es sehr wichtig, nicht nur einzelne Aspekte der Bewertung separat zu betrachten, sondern stets den Gesamteindruck einer Dolmetschleistung in ihrem jeweiligen Kontext zu bewerten, um eine umfassende Beurteilung abgeben zu können (vgl. Kutz 2005: 27).

3 Maschinelles Dolmetschen

In diesem Kapitel wird zunächst die Funktionsweise von maschinellm Dolmetschen erläutert. Dabei wird zunächst das maschinelle Dolmetschen als Prozess definiert und erklärt, ehe in den weiteren Unterkapiteln genauer auf die einzelnen Schritte einer maschinellen Dolmetschung sowie auf die dabei möglichen Problematiken beziehungsweise Schwierigkeiten eingegangen wird. Abschließend soll in diesem Kapitel der aktuelle Forschungsstand zum maschinellen Dolmetschen im Vergleich zu menschlichen Dolmetschleistungen präsentiert werden.

In dieser Arbeit wird folglich lediglich der Begriff ‚Maschinelles Dolmetschen‘ verwendet, wobei anzumerken ist, dass für diesen Prozess auch weitere Begriffe innerhalb der internationalen und interdisziplinären Forschung existieren. Vor allem in der Informatik wird oft von ‚speech-to-speech‘ gesprochen, wenn es darum geht, dass gesprochene Äußerungen von einer Software zu gesprochenen Äußerungen verarbeitet werden. Im Unterschied zu maschinellm Dolmetschen handelt es sich bei ‚speech-to-speech‘ nicht zwingend um ein Live-Setting, welches jedoch bei maschineller Dolmetschung stets gegeben ist (vgl. Fantinuoli 2025: 4f.). In der englischsprachigen Forschung wird weiters der Begriff ‚automatic speech translation‘ verwendet. Dieser steht im Kontrast zu dem Begriff ‚interpreting‘, welcher rein für menschliche Aktivitäten verwendet wird, die sich – im Gegensatz zu maschinellen Aktivitäten – durch Embodiment und Situiertheit auszeichnen. ‚Translation‘ hingegen wird auch im Englischen als Überbegriff für den Transfer einer sprachlichen Äußerung in eine andere Sprache verwendet und beinhaltet somit nicht diese spezifischen menschlichen Eigenschaften. Im Deutschen hingegen gibt es keine einfache wörtliche Übersetzung für ‚speech-to-speech‘, jedoch den im vorangehenden Kapitel erläuterten Begriff des Dolmetschens, welcher speziell mündliche Translation meint (Pöchhacker 2024: 1, 11, 17f.). Aus diesem Grund wird im Rahmen dieser deutschsprachigen Arbeit dennoch der Begriff ‚Maschinelles Dolmetschen‘ verwendet.

3.1 Funktionsweise

Maschinelles Dolmetschen bezeichnet grundsätzlich den gesamten Prozess, der aus Spracherkennung, maschineller Übersetzung und Sprachausgabe besteht. Dabei wird hier im zweiten Schritt von maschineller Übersetzung als einzelner Schritt gesprochen, da der Text, der hier von der Software bearbeitet wird, kein flüchtiger Text mehr ist, der nur einmal zur Verfügung steht, sondern es sich in diesem Schritt um eine Übersetzung in dem Sinne handelt, dass eine Translation eines schriftlichen Ausgangstextes in einen ebenso schriftlichen Zieltext erfolgt (vgl. Härtel 2016: 68).

Maschinelles Dolmetschen ähnelt folglich in gewisser Weise maschinellm Übersetzen, wobei hier noch zwei zusätzliche Schritte – die Spracherkennung und die Sprachausgabe – hinzukommen. Dies erhöht folglich die Fehleranfälligkeit einer solchen Software beziehungsweise kann zu weiteren Schwierigkeiten und Herausforderungen führen. Wie diese einzelnen Schritte genau ablaufen und welche Problematiken sich konkret ergeben können, wird im weiteren Verlauf explizit erläutert.

3.1.1 Spracherkennung

Die Spracherkennung stellt den ersten Schritt einer maschinellen Dolmetschung dar. Darunter ist der Prozess zu verstehen, im Zuge dessen ein akustisches sprachliches Signal, das von einem Mikrofon oder ähnlichem aufgezeichnet wurde, zu einer Wort- beziehungsweise Textsequenz umgewandelt wird. Die daraus entstandene Textsequenz kann anschließend für weitere sprachliche Verarbeitung, wie in diesem Fall etwa die maschinelle Übersetzung, genutzt werden (vgl. Deng & O'Shaughnessy 2003: 433). Damit eine solche Spracherkennung sich für das maschinelle Dolmetschen nutzen lässt, ist es wichtig, dass diese eine kontinuierliche Spracheingabe ermöglicht, sodass nicht jedes Wort einzeln gesprochen werden muss. Außerdem sollte die Spracherkennung sprecher:innenunabhängig funktionieren, damit diese nicht für jede:n Sprecher:in neu trainiert werden muss, um diese auch in der Praxis gut einsatzfähig machen zu können (vgl. Härtel 2016: 71f.).

Spracherkennungssoftwares arbeiten dabei auf Lexemebene, um Koartikulationseffekte, die bei kontinuierlicher Sprache zu starken Veränderungen bei der Wahrnehmung von einzelnen Phonemen führen können, möglichst ausgleichen zu können (vgl. Ruske 1994²: 5).

Dabei gibt es zwei Ansätze, die bei moderner Spracherkennung miteinander kombiniert werden (vgl. Pfister & Kaufmann 2008: 362). Statistische Sprachmodelle werden anhand von Textkorpora erstellt, wobei durch eine Analyse der Korpora Häufigkeiten berechnet und Wahrscheinlichkeiten für Wortfolgen abgeleitet werden. Dadurch erhalten auch homophone Lexeme einen Kontext und können mit einer erhöhten Trefferquote erkannt werden. Das Problem, das sich bei einem solchen Modell allerdings ergibt, liegt darin, dass es nicht möglich ist, alle akzeptablen grammatikalischen Wortfolgen abzubilden, da kein Korpus einen unbegrenzten Umfang haben kann. Grammatische Sprachmodelle hingegen arbeiten dank der Abbildung grammatischer Regeln mittels Programmiersprache wissensbasiert. Als kombinierter hybrider Ansatz ist nun statistisch gestütztes Parsing möglich, das die Wahrscheinlichkeiten für bestimmte grammatikalische Strukturen berücksichtigt. Dadurch, dass natürliche Sprache akustische Kontinuität aufweist, ist es erforderlich, dass die Spracherkennungssoftware in der Lage ist, das Gesagte automatisch zu segmentieren, was durch das Ergebnis von Wahrscheinlichkeitsberechnungen funktioniert. Außerdem können auch prosodische Informationen ausgewertet werden, um Wort- und Satzgrenzen zu erkennen (vgl. Härtel 2016: 76ff.).

Bei automatischer Spracherkennung können sich einige Problemfelder ergeben, die die Qualität der verschriftlichten Ergebnisse negativ beeinflussen können. Dazu zählen einerseits außersprachliche Störgeräusche, wie etwa Husten, Räuspern oder Trinken, die von der Software gefiltert werden müssen, um nicht fälschlicherweise als sprachliche Äußerungen erkannt zu werden (vgl. Pfister & Kaufmann 2008: 362). Außerdem stellen Selbstkorrekturen ein häufiges Phänomen natürlicher, kontinuierlich gesprochener Sprache dar, weswegen die Spracherkennungssoftware mit Pausen, Wiederholungen und Korrekturen von Wörtern umgehen können muss (vgl. Spilker et al. 2000: 131). Gesprochene Sprache besitzt außerdem im Unterschied zu geschriebener Sprache besondere sprachliche Merkmale, wie beispielsweise unvollständige Sätze oder die Verwendung von lexikalischen und grammatischen Einheiten, die nicht der schriftlichen Standardsprache entsprechen (vgl. Fantinuoli 2025: 12). Hinzu kommt, dass der Spracherkennung der non-verbale Input, wie etwa Mimik, Gestik und sonstiges Verhalten der Produzent:innen des Ausgangstextes, fehlt, welchen Humandolmetscher:innen im Idealfall wahrnehmen können. Die Analyse solchen Verhaltens kann das Verständnis des Gesagten erleichtern und Missverständnissen vorbeugen, wenn verbale Äußerungen dadurch bewusst oder unbewusst ergänzt werden (vgl. Härtel 2016: 83). Darüber hinaus ist Spracherkennung prinzipiell sprachgebunden, weswegen es zu Problemen bei der Erkennung von Fremdwörtern, Neologismen, wie zum Beispiel Neubildungen von Komposita im Deutschen, oder Eigennamen

kommen kann. Deswegen sollten diese im Idealfall im Vorfeld extra eingespeichert werden, was allerdings nicht bei jeder Software möglich ist (vgl. Härtel 2016: 73f.).

Diese Grundlagen der Funktionsweise von Spracherkennung und die möglichen Probleme, die dabei auftreten können, gilt es folglich für die spätere Analyse der maschinellen Dolmetschung zu berücksichtigen. Außerdem soll im Kapitel zur Methodik noch genauer auf die Funktionsweise der Software Google Übersetzer eingegangen werden, um zu beschreiben, welche konkreten Vor- und Nachteile die Spracherkennung dieser Software hat und welche Auswirkungen auf die Qualität des entstehenden Produktes demnach zu erwarten sind.

3.1.2 Maschinelles Übersetzen

Wie bereits erwähnt, wird der eigentliche Prozess der Translation im Zuge einer maschinellen Dolmetschung als maschinelle Übersetzung bezeichnet und gleicht grundsätzlich dem Prozess, der auch von einer Software erfolgt, deren Ein- und Ausgabeformate direkt schriftliche Texte sind.

Als maschinelles Übersetzen wird demnach jener Prozess verstanden, in dem eine spezielle Software einen Text übersetzt und somit anhand eines schriftlichen Ausgangstextes einen neuen schriftlichen Zieltext produziert. Dabei ist für die Zieltextproduktion prinzipiell kein zusätzliches menschliches Einschreiten mehr nötig. Es genügt lediglich, den Ausgangstext in die Software einzuspeisen, wodurch dieser automatisch übersetzt wird. Dies unterscheidet Software zum maschinellen Übersetzen von anderer Software wie etwa CAT-Tools, die menschliche Übersetzer:innen zwar unterstützen, wo aber trotzdem zusätzliches menschliches Handeln erforderlich ist, um einen Zieltext produzieren zu können (vgl. Forcada 2010: 215).

Dabei wird grundsätzlich zwischen zwei Ansätzen bei maschinellm Übersetzen unterschieden, welche durch ihre Funktionsweise voneinander differenziert werden können. Dazu zählen der regelbasierte und der statistikbasierte Ansatz (vgl. Forcada 2010: 218f.). Außerdem gibt es heutzutage eine Weiterentwicklung des statistikbasierten Ansatzes, bei dem zusätzlich neuronale Netzwerke verwendet werden. Dabei handelt es sich um den neuronalen Ansatz (vgl. Poibeau 2017: 82). Da für die Software Google Übersetzer der Ansatz der neuronalen maschinellen Übersetzung verwendet wird (vgl. Turovsky 2016), werden im weiteren Verlauf dieses Kapitels die beiden anderen Ansätze zwar kurz erläutert, um auf die Unterschiede aufmerksam machen zu können, der Fokus soll dabei aber insgesamt auf dem neuronalen Ansatz liegen.

Der regelbasierte Ansatz stellt den ältesten Ansatz in der Entwicklung von Software zur maschinellen Übersetzung dar. Dieser basiert dabei darauf, dass der Software sprachliche Regeln beigebracht werden, anhand derer die Strukturen des Ausgangstextes analysiert werden können und darauf aufbauend ein Zieltext generiert werden kann (vgl. Forcada 2010: 218; Härtel 2016: 85). Der Vorteil hierbei ergibt sich daraus, dass keine hohen Rechenleistungen erforderlich sind, während andererseits die Qualität der daraus entstehenden Übersetzungen meist darunter leidet, dass der Fokus rein auf grundlegenden morphologischen und syntaktischen Aspekten liegt, wodurch leicht grammatikalische Fehler entstehen können (vgl. Härtel 2016: 86).

Durch das Aufkommen von stärkerer Rechenleistung bei Computern seit den frühen 1990er Jahren wurde es möglich, große Textmengen zu sammeln und diese für die Entwicklung eines neuen Ansatzes maschineller Übersetzung zu verwenden (vgl. Poibeau 2017: 23). Bei diesem sogenannten korpusbasierten oder auch statistikbasierten Ansatz hat die Software Zugriff auf Korpora, welche anhand von auftretenden Wahrscheinlichkeiten analysiert werden. Das bedeutet, dass ein gewisser statistisch messbarer Zusammenhang zwischen Strukturmerkmalen des Ausgangs- und jenen des Zieltextes besteht. So kann beispielsweise ausgerechnet werden, mit welcher Wahrscheinlichkeit ein bestimmtes Wort auf ein anderes folgt oder mit welcher Wahrscheinlichkeit ein Wort die Übersetzung eines anderssprachigen Wortes darstellt. Außerdem können weitere Parameter wie etwa die Satzlänge oder die Position von Wörtern in Sätzen in diese Berechnungen miteinfließen. Weiters können einsprachige Texte als zusätzlicher Input verwendet werden, wodurch die linguistischen Strukturen der jeweiligen Sprache noch besser von solcher Software erkannt werden können (vgl. Härtel 2016: 88f.).

Dank der Verwendung von neuronalen Netzen bei dem neuronalen Ansatz ist solche Software auch in der Lage, komplexere Konzepte auf Basis von wenigen Informationen zu kreieren. Das bedeutet, dass nur mehr wenige Elemente manuell spezifiziert werden müssen und das System daraus aufgrund von Deep Learning selbst neue Daten ableiten kann (vgl. Poibeau 2017: 82f.). Das bedeutet weiters, dass für den gesamten Übersetzungsprozess durch eine solche Software ein einzelnes, großes neuronales Netzwerk zum Einsatz kommt, anstatt lediglich Korpora zu verwenden. Dadurch ist es außerdem möglich, dass die daraus resultierenden Zieltexte meist als flüssiger zu lesen wahrgenommen werden als dies bei Übersetzungen durch den statistikbasierten Ansatz der Fall ist (vgl. Zhang et al. 2020: 21). Hinzu kommt, dass bei diesem Ansatz nicht nur ganze Sätze als solche betrachtet werden, sondern ebenso der Kontext – also die umliegenden Sätze – in die Übersetzung miteinbezogen werden können (vgl. Poibeau 2020: 84).

Auch bei maschineller Übersetzung gibt es Kriterien, anhand derer die Qualität der produzierten Zieltexte gemessen wird. Da sich diese Arbeit allerdings mit dem gesamten Prozess der maschinellen Dolmetschung befasst, wird sich die spätere Analyse auf die Kriterien, die in dem vorangegangenen Kapitel erarbeitet wurden, stützen. Nichtsdestotrotz sollen an dieser Stelle jene Kriterien kurz genannt werden, die in der Translationswissenschaft in Zusammenhang mit maschineller Übersetzung stehen. Dazu zählen Flüssigkeit, Genauigkeit, Fehlerlosigkeit, Informationsgehalt und Angemessenheit (vgl. Hansen 2010: 387; Koehn 2020: 4).

Auf Basis der Funktionsweise von maschineller Übersetzung und den geforderten Kriterien lassen sich auch hier Probleme beziehungsweise Schwierigkeiten feststellen, die in Zusammenhang mit maschineller Übersetzung auftreten können. Diese sind wiederum wichtig zu beachten, um sie in der späteren Analyse auch erkennen und erklären zu können.

Zu solchen Schwierigkeiten zählt beispielsweise, dass einer Software für maschinelles Übersetzen situatives Wissen fehlt. Das bedeutet, dass ein Zieltext nicht an spezifische Gegebenheiten der jeweiligen Situation angepasst werden kann. Konkret gesagt betrifft dies etwa die Verwendung von Höflichkeitsformen. Diese ist dabei zusätzlich stark sprachenspezifisch, da je nach Sprachkombination teilweise ähnliche Höflichkeitsformen verwendet werden können, während in anderen Fällen die jeweiligen Formen erst angepasst werden müssen (vgl. Härtel 2016: 90). Dies betrifft dabei vor allem etwa Sprachkombinationen, bei welchen eine Sprache beispielsweise gar keine Unterscheidung zwischen verschiedenen Pronomen für die direkte Ansprache vornimmt, während dies in der anderen wichtig wäre. Ein Beispiel hierfür wäre die Sprachkombination Englisch-Deutsch. Da sich diese Arbeit jedoch mit der Sprachkombination Französisch-Deutsch befassen wird, sollte dieser Punkt kein allzu großes Problem darstellen, da die Verwendung der Pronomen für die direkte Ansprache in diesen beiden Sprachen ähnlich funktioniert. Es liegt nicht nur eine eindeutige Unterscheidung zwischen den unterschiedlichen Pronomen für die direkte Ansprache vor, sondern auch die pragmatische Verwendung der Höflichkeitsformen in Ausgangs- und Zielkultur sehr ähnlich sind, sodass hierbei keine groben Fehler zu erwarten sind.

Zusätzlich zur Anrede können sich weitere Schwierigkeiten bei der Übersetzung in Bezug auf kulturelle Unterscheide ergeben. Dies lässt sich damit erklären, dass es für eine Software für maschinelle Übersetzung meist nicht möglich ist, außersprachliches Wissen in die Übersetzung miteinzubeziehen. Dadurch können sich bei der Übersetzung von beispielsweise Sprichwörtern, kulturellen Referenzen, Redewendungen, Sarkasmus oder Ironie Probleme ergeben, da in solchen Fällen oft keine direkte Übersetzung möglich ist, sondern kontextuelles und außersprachliches Wissen in die Übersetzung miteinbezogen werden muss. Demnach sollte

eine solche Software idealerweise derartige Informationen erkennen können, um diese in äquivalente Ausdrücke in der Zielsprache zu übertragen (vgl. Fantinuoli 2025: 13f.).

Eine weitere Schwierigkeit können transphrastische Bezüge darstellen. Dies bezieht sich etwa auf Anaphern oder Kataphern, also Ausdrücke, die sich auf frühere oder spätere Textstellen – auch über Satzgrenzen hinweg – beziehen. Diese Verknüpfungen innerhalb eines Textes sind wichtig, um einen in sich logisch aufgebauten und kohärenten Text zu erstellen, der von den Rezipient:innen problemlos verstanden werden kann (vgl. Härtel 2016: 89f.). Dies ist vor allem dann relevant, wenn die Software auf Satzebene segmentiert und darüber hinaus keinen Kontext berücksichtigt. Die heute verwendete Software nach dem neuronalen Ansatz hingegen sollte – wie bereits oben erwähnt – auch den Kontext miteinbeziehen, weswegen es abzuwarten gilt, ob diese Schwierigkeit auch auf die Software Google Übersetzer zutreffen wird.

Hinzu kommt ein weiteres großes Problemfeld, das sich bei maschineller Übersetzung ergeben kann. Dies betrifft die Ambiguität von Sprache. Diese Ambiguität kann sich auf verschiedene Ebenen des Sprachsystems beziehen. Das bedeutet, dass sowohl die Syntax als auch die Lexik beziehungsweise die Semantik davon betroffen sein können. Ambiguität meint dabei, dass eine sprachliche Form verschiedene Bedeutungen haben kann, welche sich meistens aus dem situativen oder textuellen Kontext ergeben (vgl. Forcada 2010: 215). Ein Beispiel für lexikalische Ambiguität im Deutschen wäre etwa der Homograph *umfahren*, welcher je nach Betonung in der gesprochenen Sprache eine andere Bedeutung haben kann. In schriftlichen Texten jedoch muss diese Bedeutungsunterscheidung rein aus dem Kontext getroffen werden, was wiederum für maschinelle Übersetzungssoftware ein Problem darstellen könnte. Syntaktische Ambiguität bezieht sich auf ganze Sätze, die unterschiedliche Bedeutungen haben können. Ein bekanntes Beispiel aus dem Deutschen wäre etwa der Satz *Ich sehe den Mann mit dem Fernrohr*, wobei hier ohne Kontext unklar ist, ob der Mann oder das Ich-Subjekt in diesem Satz das Fernrohr besitzen.

Ein weiteres Problem könnten Unterschiede im Satzbau der jeweiligen Sprachen darstellen. Dies betrifft dabei vor allem Sprachen, die strukturell sehr unterschiedlich zueinander sind, da hierbei mehr Änderungen bezüglich der syntaktischen Struktur seitens der Software erfolgen müssen, wodurch sich wiederum mehr Fehlerquellen ergeben können (vgl. Forcada 2010: 216). In Bezug auf die Untersuchung in der Sprachrichtung Französisch-Deutsch sollte dieser Punkt voraussichtlich zu keinen größeren Problemen führen, da es im Französischen eine festgelegte Reihenfolge bezüglich der Anordnung der Satzglieder gibt, welche im Deutschen meistens ebenso übernommen werden kann, da das Deutsche eine recht freie syntaktische Anordnung erlaubt.

Bei der Bewertung von Texten, die durch eine Software für das maschinelle Übersetzen produziert wurden, ist es folglich wichtig, vor allem sprachliche Ambiguitäten sowie deiktische Elemente zu betrachten, da es hier ausschlaggebend ist, sowohl den textuellen als auch den situativen Kontext der Ausgangstextproduktion zu berücksichtigen, um einen inhaltlich korrekten und kohärenten Zieltext produzieren zu können. In der späteren Analyse und dem Vergleich der Dolmetschleistungen zwischen den Humandolmetscher:innen und der Software soll also ebenso erarbeitet werden, ob diese Problematiken bei Google Übersetzer auftreten und inwieweit dadurch die Qualität der Dolmetschleistung leidet.

3.1.3 Sprachausgabe

Die Sprachausgabe stellt den finalen Schritt einer maschinellen Dolmetschung dar und sorgt dafür, dass die Rezipient:innen den Zieltext, der im Zuge der maschinellen Übersetzung produziert wurde, auch akustisch wahrnehmen können.

Damit eine solche Sprachausgabe funktionieren kann, muss der schriftsprachliche Text zunächst in eine eindeutige, strukturierte Repräsentation umgewandelt werden. Im Zuge dieser Repräsentation wird der Text nach Sätzen, Lexemen und Syngraphemen segmentiert, was die Grundlage für die anschließende Sprachausgabe darstellt (vgl. Härtel 2016: 93). Dabei können Homographie, deren semantische Bedeutung von der Aussprache oder der Betonung abhängt, ein potentiell Problem darstellen (vgl. Härtel 2016: 97), weswegen es wichtig ist, dass die einzelnen Segmente für eine erfolgreiche Sprachausgabe anhand ihrer syntaktischen Funktion klassifiziert werden. Demnach gilt es für die Software, im Rahmen der oben genannten Repräsentation des schriftlichen Textes auch eine Textanalyse durchzuführen, um eine Disambiguierung von homographen Lexemen zu gewährleisten (vgl. Taylor 2009: 64). Ähnliches gilt auch für syntaktische Ambiguität, deren unterschiedliche Bedeutungen etwa ebenso von der Betonung der jeweiligen syntaktischen Einheiten abhängen können (vgl. Härtel 2016: 97). Hinzu kommt, dass es spezifischer Ausspracheregeln für Eigennamen, Fremdwörter, Abkürzungen oder Akronyme bedarf. Dies kann insbesondere zu Problemen führen, wenn es sich dabei um fremdsprachliche Ausdrücke handelt, die nicht den gewohnten Ausspracheregeln der jeweiligen Sprache folgen. Damit also die lautsprachliche Synthese solcher Ausdrücke bei der Sprachausgabe gelingt, müssen solche Ausnahmen im Zweifelsfall manuell in der Software hinterlegt

werden, wenn nicht davon ausgegangen werden kann, dass die korrekte Aussprache anhand der bereits programmierten Regeln realisiert werden kann (vgl. Taylor 2009: 209).

Neben der grundsätzlichen Anforderung an eine korrekte Aussprache gibt es seitens der Nutzer:innen von Sprachausgabe auch die Erwartung, dass eine synthetische Stimme wie eine menschliche Stimme klingen sollte, obwohl auch bei einer synthetisch wirkenden Stimme die Verständlichkeit gleichermaßen gut gegeben sein kann. Dabei spielt ebenso die Prosodie eine wichtige Rolle für die Rezipient:innen, worauf in Kapitel 2.2. bereits genauer eingegangen wurde. Dementsprechend würde es den Erwartungen der Nutzer:innen entsprechen, wenn die Software der Sprachausgabe auch Prosodie generieren kann. Dabei ergibt sich jedoch das Problem, dass die Schriftsprache nur stark eingeschränkt Hinweise auf prosodische Elemente enthält. Diese würden zwar theoretisch durch eine Analyse des Ausgangstextes vorliegen, jedoch muss die Software dabei in der Lage sein, diese Informationen nicht nur zu analysieren, sondern auch in die Sprachausgabe miteinarbeiten zu können (vgl. Härtel 2016: 98f.).

Bei der späteren Bewertung der Sprachausgabe von Google Übersetzer soll speziell auf die hier genannten möglichen Problematiken der korrekten Aussprache von beispielsweise Homographen, Fremdwörtern oder Eigennamen geachtet, sowie analysiert werden, inwiefern sich der Klang der Stimme sowie die Prosodie auf die Qualität der Sprachausgabe auswirken.

3.2 Forschungsstand

In diesem Kapitel soll auf den aktuellen Forschungsstand bezüglich des maschinellen Dolmetschens beziehungsweise den Vergleich zwischen Dolmetschleistungen einer Software und Dolmetscher:innen eingegangen werden. So können in weiterer Folge die Präsentation der Methodik sowie die eigentliche Analyse und der Vergleich der maschinellen mit der humanen Dolmetschung in den bisherigen Forschungskontext eingebettet und deren Relevanz hervorgehoben werden.

Hierbei sollen vor allem Arbeiten erwähnt werden, die in den letzten Jahren ebenfalls das Thema der maschinellen Dolmetschung behandelt haben. Sarah Fünfer (2013) arbeitete in ihrer Untersuchung mit dem Programm ‚Lecture Translator‘, welches am Karlsruher Institut für Technologie eigens für die Dolmetschung von Vorlesungen konzipiert wurde. So wurden in dieser Untersuchung Simultandolmetschungen der Software und der Humandolmetscher:innen aus dem Englischen ins Deutsche, ins Französische und ins Spanische miteinander

verglichen. Anschließend erfolgte nicht nur eine Befragung der Zuhörer:innen nach ihren Einschätzungen, sondern auch ein analytischer Vergleich der Dolmetschungen anhand von Qualitätskriterien. Diese Qualitätskriterien glichen dabei stark jenen, welche auch im Rahmen dieser Arbeit verwendet werden sollen. Fünfer (2013) kam schließlich zu dem Ergebnis, dass die maschinellen Dolmetschleistungen große Kohärenzprobleme aufwiesen. Dazu kamen einige unvollständige Sätze sowie diverse sprachliche Fehler. Das Problem hierbei ist, dass es dadurch auch zu inhaltlichen Missverständnissen kommen kann. Im Gegensatz dazu gab es keine Auslassungen, welche dafür teilweise bei den Humandolmetscher:innen auftraten. Die Befragung der Zuhörer:innen bestätigte dieses Ergebnis, weshalb klar wurde, dass diese Software nicht mit den Leistungen von Humandolmetscher:innen mithalten kann (vgl. Fünfer 2013: 55, 80, 125-128).

Fantinuoli und Prandi (2021) verglichen die Leistung von einem:r menschlichen Dolmetscher:in mit jener einer Software für Speech-to-text-Translation, also einem System, das gesprochene Sprache verschriftlicht, übersetzt und für Rezipient:innen in Echtzeit einen geschriebenen Text zur Verfügung stellt. Dabei wurde mit dem Sprachenpaar Englisch-Italienisch gearbeitet und der Fokus auf die zwei Qualitätsparameter Verständlichkeit und Informationsgehalt gelegt. Als Ergebnis wird festgehalten, dass die Verständlichkeit bei den menschlichen Dolmetschleistungen deutlich höher als bei jener der Software liegt, während der maschinelle Output in Bezug auf den Informationsgehalt etwas besser als jener des Menschen zu bewerten ist. Dies ist damit zu begründen, dass bei menschlichen Dolmetschleistungen Generalisierungen und bewusste Auslassungen erfolgen, die bei maschineller Dolmetschung oder Übersetzung nicht vorkommen. Bei Kombination beider Kriterien erreicht der:die Humandolmetscher:in eine deutlich höhere Bewertung, wobei angemerkt wird, dass ein Vergleich mit mehreren menschlichen Leistungen bei einer größeren Vielfalt an Texten sinnvoll wäre, da bei menschlichen Dolmetschungen eine gewisse Varianz besteht (vgl. Fantinuoli & Prandi 2021: 247, 249f.).

Außerdem sollen hier drei Masterarbeiten erwähnt werden. Dazu zählt die Arbeit von Giulia Cürten (2016), welche sich mit Dolmetschleistungen von ‚Google Übersetzer‘² in einem dialogischen Setting – einer Situation in einem Krankenhaus und einem Dialog mit Tourist:innen – befasste. Es wurden vorgeschriebene Dialoge im Sprachenpaar Deutsch-Italienisch, die jeweils von unterschiedlichen Stimmen eingesprochen wurden, untersucht. Dabei konnte

² Die Version von ‚Google Übersetzer‘, die Cürten (2016) untersuchte, besaß noch nicht jene Transkriptionsfunktion, welche in der vorliegenden Masterarbeit untersucht wird. Darum ergaben sich hier auch die zu kurzen Verarbeitungskapazitäten, die in der heutigen Version keine Rolle mehr spielen.

einerseits festgestellt werden, dass die Verarbeitungskapazität der Spracherkennung dieser Software für einige Dialogpassagen zu gering war, weswegen zusätzliche Pausen eingefügt werden mussten. Andererseits erkannte die Spracherkennung alle Stimmen gleichermaßen und auch einige wichtige Termini, die in diesen Kontexten vorkommen, wurden korrekt gedolmetscht. Insgesamt kam Cürten (2016) allerdings zu einem ähnlichen Ergebnis wie auch Fünfer (2013). Die Dolmetschungen wiesen einige sprachliche sowie grammatikalische Fehler auf, die zu inhaltlichen Missverständnissen führen können. Außerdem wurde die fehlende Intonation der Zieltextproduktion bemängelt sowie einige Probleme bei der Spracherkennung festgestellt, welche wiederum zu Translationsfehlern führten. Auch hier wurde das Ergebnis bestätigt, dass diese Software nicht mit menschlichen Dolmetscher:innen mithalten kann, sofern die Ausgangstexte über sehr simple Alltagsgespräche hinausgehen (vgl. Cürten 2016: 23f., 114ff.).

Julia Leitner (2020) befasste sich in ihrer Untersuchung mit einem Vergleich zwischen Dolmetschungen der Software ‚iTranslate‘ und einer Humandolmetschung. Auch diese Dolmetschungen waren dabei in ein dialogisches Setting – ein Polizeigespräch – eingebettet. Es handelte sich dabei um Dialoge im Sprachenpaar Französisch-Deutsch, die spontan anhand eines Gesprächsleitfadens produziert wurden. Auch Leitner (2020) kam anhand des Vergleichs zwischen der menschlichen und der maschinellen Dolmetschung zu einem ähnlichen Ergebnis. Bei der maschinellen Dolmetschung wurden ebenfalls viele Fehlerquellen ausgemacht, zu denen etwa Grammatik, Lexik oder auch Syntax zählen, welche wiederum zu inhaltlichen Missverständnissen führen können. So lautete auch hier das Fazit, dass eine solche Software für den kurzen alltäglichen Gebrauch durchaus zu verwenden ist und akzeptable Ergebnisse liefern kann, sie allerdings nicht mit der Leistung von professionellen Kommunal Dolmetscher:innen mithalten kann und somit in einem solchen Setting weniger gut funktioniert (vgl. Leitner 2020: 55f., 112).

Julia Glocknitzer (2020) wiederum untersuchte die Leistung der Dolmetsch-App ‚SayHi Translate‘ im Rahmen eines Arzt-Patient-Gesprächs mit dem Sprachenpaar Deutsch-Ungarisch. Auch hier handelt es sich folglich um ein dialogisches Setting, in welchem der Patient rein auf die Hilfe der App angewiesen war, um verbal kommunizieren zu können. Glocknitzer (2020) kam – wie auch bereits in vorherigen Arbeiten anhand anderer Apps gezeigt wurde – zu dem Ergebnis, dass die Funktionsweise der untersuchten App nach ihrem derzeitigen Stand nicht ausreicht, um eine einwandfreie Kommunikation im medizinischen Bereich zu gewährleisten, weswegen die Leistung von professionellen Dolmetscher:innen im Gesundheitswesen dadurch nicht ersetzt werden kann. Dies lässt sich sowohl mit Fehlern bei der Spracherkennung als auch mit Übersetzungsfehlern begründen, die in weiterer Folge ganze Textpassagen

unverständlich machen oder den Inhalt des Ausgangstextes falsch wiedergeben (vgl. Glocknitzer 2020: 112f, 115).

In der vorliegenden Masterarbeit soll aufbauend auf den aktuell vorliegenden Ergebnissen ein weiterer Aspekt von maschineller Dolmetschung untersucht werden, der bisher noch nicht behandelt werden konnte. Dabei handelt es sich um eine Konsektivdolmetschung innerhalb eines Konferenzsettings, die zuvor noch nicht für ähnliche Untersuchungen möglich war, da es bei den bisher untersuchten Apps nicht möglich war, längere zusammenhängende Textabschnitte per Spracherkennung am Stück aufzuzeichnen, um diese anschließend dolmetschen zu können. Dementsprechend soll in dieser Arbeit nun die Möglichkeit genutzt werden, eine gesamte mehrminütige abgeschlossene Rede konsektiv zu dolmetschen. So soll herausgefunden werden, ob dabei ähnliche Probleme wie die oben beschriebenen auftreten, oder ob die aktuell verfügbare Software mittlerweile schon derart weiterentwickelt werden konnte, um die Qualität einer maschinellen Dolmetschung mit einer menschlichen Dolmetschung vergleichbar und somit auch in der Praxis einsatzfähig zu machen. Zusätzlich werden mehrere menschliche Dolmetschleistungen mit jener der Software verglichen, um die eventuelle Varianz zwischen den einzelnen menschlichen Leistungen ebenfalls in die Analyse miteinbeziehen zu können und ein allgemeineres Bild dieser Dolmetschungen zu erhalten, worauf im folgenden Kapitel noch genauer eingegangen wird.

4 Methode

Im Zuge der vorliegenden Arbeit soll die eingangs erwähnte Forschungsfrage nach den Möglichkeiten und Grenzen der Software ‚Google Übersetzer‘ im Vergleich zu menschlichen Dolmetscher:innen beantwortet werden. Dafür soll in diesem Kapitel die Methodik vorgestellt werden, nach der die folgende Analyse durchgeführt wurde. Zuerst erfolgt die Erläuterung der Datenerhebung. Diese besteht aus der Darlegung des Ausgangstextes und der anschließenden Erläuterung der Vorgangsweise bei der maschinellen Dolmetschung sowie bei den Humandolmetschungen. Daraufgehend wird erklärt, wie die so erhobenen Daten für die abschließende Analyse transkribiert wurden, ehe in einem zweiten Unterkapitel dargelegt wird, inwiefern die Datenanalyse ablaufen soll.

4.1 Datenerhebung

Die Datenerhebung lief wie folgt ab: Als Ausgangstext für die Dolmetschung wurde die französischsprachige Rede „La reconnaissance faciale“³ aus der Datenbank ‚Speech repository‘ der Europäischen Kommission ausgewählt, die ins Deutsche gedolmetscht werden soll. Diese Rede dauert etwa sechs Minuten und ist – laut den Angaben der Datenbank – für eine Konsekutivdolmetschung auf einem hohen Niveau konzipiert und kann demnach in Bezug auf den Schwierigkeitsgrad und die Dauer der Rede mit realen Einsätzen verglichen werden. Der Vorteil bei dieser Rede liegt darin, dass die Audioqualität sehr gut ist und die Sprecherin eine deutliche Sprechweise und ein angenehmes Sprechtempo aufweist. Hinzu kommt, dass es keine störenden Hintergrundgeräusche gibt und die Rednerin ebenso gut vorbereitet ist und dementsprechend wenig Versprecher, Verzögerungslaute und Selbstkorrekturen produziert. Die Rede liegt außerdem zwar im Videoformat vor, jedoch spielt der nonverbale Input, den menschliche Dolmetscher:innen wahrnehmen könnten, keine allzu große Rolle in Bezug auf das Verständnis des Gesagten. So verwendet die Sprecherin etwa keine starke Gestik oder Mimik und es werden auch keine zusätzlichen visuellen Informationen etwa durch Einblendungen von Bildern, Texten oder Ähnlichem dargeboten. Demnach weist der Ausgangstext wenig potentielle Störfaktoren für eine maschinelle Dolmetschung auf, weswegen es sich klarerweise um ein etwas

³ <https://speech-repository.webcloud.ec.europa.eu/speech/la-reconnaissance-faciale> (Stand: 24.05.2026)

idealisiertes Setting handelt. Jedoch soll damit ausgeschlossen werden, dass die maschinelle Dolmetschung zu stark von außersprachlichen Faktoren negativ beeinflusst wird, damit sich die Analyse vor allem auf die sprachlichen Aspekte konzentrieren kann.

Da eine Simultandolmetschung mit der App Google Übersetzer ohne weitere technische Hilfsmittel noch nicht möglich ist, wird die Dolmetschung im Konsekutivmodus durchgeführt. Dabei sieht der Ablauf folgendermaßen aus: Innerhalb der App Google Übersetzer auf einem Smartphone werden zunächst die Ausgangssprache Französisch und die Zielsprache Deutsch in den entsprechenden Feldern ausgewählt. Weitere Eingaben, wie etwa das Trainieren der Spracherkennungssoftware oder die zusätzliche Eingabe von Fachtermini, Hinweisen zur Aussprache von Fremdwörtern, Eigennamen oder Ähnlichem sind nicht möglich. Um die maschinelle Dolmetschung zu beginnen, wird die Transkriptionsfunktion durch das Tippen auf das Mikrofonsymbol gestartet. Anschließend wird das Video der Rede von einem zweiten Gerät über Lautsprecher abgespielt und von der App mithilfe der Spracherkennungsfunktion aufgenommen. Diese produziert in Echtzeit ein Transkript des Gesagten, welches entweder auf Französisch oder auch gleich als deutsche Übersetzung auf dem Bildschirm des Smartphones angezeigt werden kann. Dies hat den Vorteil, dass sowohl das eigentliche ausgangssprachliche Transkript des Gesagten als auch das Ergebnis der maschinellen Übersetzung in der Zielsprache, welche sich beide in voller Länge im Anhang finden, einzeln analysiert werden können, um zu eruieren, welche Probleme oder potentiellen Fehler bereits bei der Spracheingabe und welche erst bei der maschinellen Übersetzung aufgetreten sind. Nach dem Ende der Rede kann durch Klicken auf das entsprechende Symbol die bereits vorliegende deutsche maschinelle Übersetzung des erstellten französischen Transkriptes dank der Sprachausgabefunktion der App Google Übersetzer über die Lautsprecher des Smartphones lautsprachlich wiedergegeben werden. Dadurch ergibt sich aus der Spracherkennung, der darauffolgenden Transkription, der maschinellen Übersetzung und der abschließenden Sprachausgabe eine vollständige Konsekutivdolmetschung.

Als Vergleichsbasis für die Bewertung der maschinellen Dolmetschung werden die Konsekutivdolmetschungen von vier Dolmetscher:innen herangezogen. Diese befinden sich alle im Masterstudium Translation mit dem Schwerpunkt Konferenzdolmetschen an der Universität Wien. Alle Dolmetscher:innen haben Französisch und Deutsch als Arbeitssprachen, wobei eine davon – im Gegensatz zu den anderen drei – Deutsch nicht als Erstsprache hat. Zusätzlich haben alle bereits Lehrveranstaltungen zum Konsekutivdolmetschen besucht beziehungsweise teilweise die abschließende Dolmetschprüfung des Studiums absolviert und sind somit ausreichend qualifiziert, um als Vergleichsbasis dienen zu können.

Die Rahmenbedingungen für die Dolmetscher:innen waren dabei wie folgt: Die Dolmetscher:innen meldeten sich nach einem Aufruf mit Bekanntgabe der Bedingung, dass sie sich in Endphase des Masterstudiums Konferenzdolmetschen mit den Arbeitssprachen Deutsch und Französisch befinden sollten, freiwillig, um an diesem Experiment teilzunehmen. Da die Dolmetschungen zwischen Juli 2021 und Oktober 2022 und somit in Zeiten der Kontaktbeschränkungen aufgrund der Coronapandemie durchgeführt wurden, fand der Kontakt mit den Dolmetscher:innen rein digital statt.

Die Dolmetscher:innen erhielten von mir folgende Instruktionen: Es soll keine explizite Vorbereitung auf das Thema erfolgen, wie beispielweise Recherchen zum Thema der Rede, die Dolmetschung anderer Reden zum gleichen Thema oder das Erstellen von Glossaren oder Ähnlichem. Demnach hatten die Dolmetscher:innen nur jene Zusatzinformationen zur Verfügung, die auf der Internetseite des Speech Repository, wo auch die eigentliche Rede aufscheint, zu finden sind. Bei diesen Informationen handelt es sich um den Titel der Rede „La reconnaissance faciale“ [Die Gesichtserkennung], die Dauer der Rede, die mit sechs Minuten und fünfzehn Sekunden angegeben ist, die – oben bereits erwähnte – hohe Schwierigkeit der Rede sowie eine kurze Liste an Fachtermini, die in der Rede vorkommen. Hierbei handelt es sich um folgende Fachbegriffe: s'inmiscer [sich einschleichen], friande [begierig], intrusif [aufdringlich], enregistrement sous identité réelle [Registrierung unter Angabe der echten Identität], données biométriques [biometrische Daten], empreintes digitales [Fingerabdruck], scans de l'iris de l'œil [Iris-Scans]. Demnach durften die Dolmetscher:innen lediglich die hier aufgeführten Termini vorab nachschlagen.

Die Dolmetschungen selbst wurden von den Dolmetscher:innen aufgrund der bereits erwähnten Kontaktbeschränkungen an ihrem jeweiligen Aufenthaltsort selbstständig durchgeführt. So hatten die Dolmetscher:innen vorab kurz die Möglichkeit zur Recherche der Fachtermini, ehe sie sich die Rede einmal anhören konnten. Während dem Zuhören war es erlaubt, händische Notizen anzufertigen, wie es beim Konsekutivdolmetschen üblich ist. Abgesehen davon waren keinerlei weitere Hilfsmittel, wie Wörterbücher, Transkriptionsprogramme oder Ähnliches erlaubt. Nach dem Anhören der Rede sollte dann die Wiedergabe in der Zielsprache Deutsch erfolgen, wobei die Dolmetscher:innen dabei jeweils selbst eine Audioaufnahme ihrer Darbietung des Zieltextes anfertigten. Diese Audioaufnahmen wurden mir übermittelt und anschließend von mir transkribiert. Diese Transkripte sind in vollem Umfang im Anhang der vorliegenden Arbeit zu finden.

Die Konventionen bei den von mir durchgeführten Transkriptionen sind grundsätzlich an Rehbein (2004) angelehnt, wurden aber von mir selbst speziell für die Bedingungen hier

angepasst. Demnach erfolgte die Erstellung der Transkripte nach den folgenden Konventionen: Als Format für die Transkriptionen wurde aufgrund der einfachen Lesbarkeit ein Fließtext gewählt, der mit Zeilennummern zur besseren Nachvollziehbarkeit der Analyse ausgestattet wurde. Bei der Verschriftlichung des Gesprochenen erfolgte außerdem die Verwendung der Standardorthographie des Deutschen, was in diesem Fall passend war, da keine dialektalen Ausdrücke in den Dolmetschungen vorkamen. Dies gilt ebenfalls für die Verwendung von Satzzeichen, wie etwa Punkte, Kommata und Fragezeichen. Auch diese wurden nach den Regeln der deutschen Orthographie gesetzt.

Wichtig für die spätere Analyse der Dolmetschungen waren ebenso die Markierungen von Pausen, Betonungen, Selbstkorrekturen, Verzögerungslauten und anderen möglichen Störgeräuschen, da dadurch auch Rückschlüsse auf die Prosodie des Gesagten gezogen werden können. Demnach erfolgte die Markierung von Pausen durch das Zeichen • für sehr kurze Intonationspausen, durch das Zeichen • • von kurzen Pausen im Redefluss und durch die Angabe von beispielsweise ((3s)) als Zeichen für längere Pausen ab zwei Sekunden mit genauer Sekundenangaben. Die Angabe von Betonungen erfolgte durch die Großschreibung von auffällig betonten Silben, wobei normale Wort- und Satzbetonungen nicht explizit gekennzeichnet wurden, um die Lesbarkeit der Transkripte nicht zu beeinträchtigen und hervorstechende Betonungen deutlich erkennbar zu machen. Auffällige Dehnungen von Lauten wurden durch das Setzen eines Doppelpunktes nach dem betreffenden Laut markiert, wobei Ähnliches wie bei den Betonungen gilt, sodass nur hervorstehende Dehnungen markiert wurden. Selbstkorrekturen und abgebrochene Worte wurden an den jeweiligen Stellen durch das Setzen eines Schrägstrichs gekennzeichnet. Deutlich hörbares Ein- und Ausatmen der Dolmetscher:innen wurde durch ((ea)) beziehungsweise ((aa)) gekennzeichnet. Verbale Verzögerungslaute, wie beispielsweise ‚äh‘ wurden als solche direkt verschriftlich, während sonstige außersprachliche Störgeräusche wie Räuspern, Husten oder das Rascheln von Papier in doppelter Klammer, wie beispielsweise ((Räuspern)) verschriftlicht wurden.

Diese von mir erarbeiteten Transkriptionskonventionen wurden auch bei der Transkription des französischsprachigen Ausgangstextes, sowie bei der Überarbeitung des von Google Übersetzer erstellten Transkriptes des deutschen Zieltextes angewandt, um diesen ebenfalls prosodisch analysieren und eine bessere Vergleichbarkeit zwischen der maschinellen und den menschlichen Dolmetschungen gewährleisten zu können. Zusätzlich werden bei der maschinellen Dolmetschung die einzelnen Schritte anhand der einzelnen Transkripte getrennt voneinander analysiert, um herauszufinden, welcher der drei Schritte – Spracheingabe, maschinelle

Übersetzung, Sprachausgabe – welche Problematiken mit sich bringt. Die genauen Kriterien und die Vorgehensweise bei dieser Datenanalyse werden im folgenden Kapitel erläutert.

4.2 Datenanalyse

Die Bewertung und Einordnung der einzelnen Dolmetschleistungen und die Gegenüberstellung zwischen der maschinellen und den menschlichen Dolmetschungen erfolgte anhand der im Rahmen von Kapitel 2 genauer beschriebenen Kriterien, den Qualitätsanforderungen und des für die Analyse erarbeiteten Bewertungsschemas, welches aus einer qualitativen und einer quantitativen Analyse besteht und zwei Bewertungsgrundlagen vorsieht. Dabei handelt es sich einerseits um die Übertragung des Inhalts und andererseits um die Realisierung der Dolmetschung. Bezüglich der Übertragung des Inhalts gilt es, eine inhaltliche Treue zum Ausgangstext zu bewahren, sodass es zu keinen ‚information gains‘ oder ‚information losses‘ oder sogar zu Widersprüchen zwischen Ausgangs- und Zieltext kommt. Weiters sollten pragmatische Besonderheiten angemessen übertragen werden und der produzierte Zieltext sollte ausreichend Kohäsion und Kohärenz aufweisen. Zu der Realisierung der Dolmetschung zählen zum einen stimmliche Faktoren wie der Klang der Stimme, das Sprechtempo, der Rhythmus, die Intonation und die Aussprache und zum anderen die Sprachkompetenz. Darunter ist zu verstehen, dass Grammatik, Syntax, Lexik und Stil möglichst fehlerlos und dem Kontext und Text angemessen verwendet werden. Zusätzlich gilt es noch, die erwartete Dauer einer Konsekutivdolmetschung von etwa 70% des Originals einzuhalten.

Neben dieser qualitativen Analyse der Dolmetschleistungen und dem Vergleich zwischen Mensch und Maschine wurde zusätzlich eine kurze quantitative Fehleranalyse durchgeführt, um eine noch bessere und umfangreichere Übersicht über die Unterschiede in Hinblick auf die Qualität zwischen den einzelnen Dolmetschungen zu erhalten. Damit diese quantitative Analyse adäquat in Bezug zu der qualitativen Analyse gesetzt werden kann, kommen hierfür grundsätzlich die gleichen und oben erläuterten Bewertungskriterien zum Einsatz, wobei nur jene in Betracht kommen, die auch quantitativ ausgewertet werden können. So werden einerseits Fehler in Bezug auf die Übertragung des Inhalts und damit Fehler bei der Treue zum Ausgangstext, der Kohärenz beziehungsweise Kohäsion und der Pragmatik und andererseits Fehler bei der Realisierung der Dolmetschung gewertet. Dabei zählen Fehler bezüglich der Grammatik, der Lexik und der Syntax.

Diese in der qualitativen Analyse erkannten und beschriebenen Fehler aus den erwähnten Kategorien werden in weiterer Folge nach ihrem Schweregrad in drei Fehlerklassen – leichte, mittlere und schwere Fehler – einsortiert und anschließend gezählt. Der Schweregrad bezieht sich darauf, in welchem Maße die geforderten Qualitätskriterien umgesetzt wurden, sodass ein jeweils leichter, mittlerer oder schwerer Verstoß für den entsprechenden Fehlertyp stehen soll. Ein leichter Fehler bedeutet somit etwa, dass zwar ein Fehler durch die qualitative Analyse erkennbar ist, dieser aber das Verständnis seitens des Zielpublikums nicht wirklich einschränkt, wenn es sich beispielsweise um einen lexikalischen Fehler handelt oder für das Zielpublikum gar nicht oder nur kaum relevant ist, wenn es beispielsweise die Auslassung von Inhalt betrifft. Ein mittlerer Fehler bedeutet folglich, dass das Verständnis stärker beeinträchtigt wird, beziehungsweise dass inhaltliche Einbußen zu erkennen sind. Unter einem schwereren Fehler ist zu verstehen, dass das Verständnis stark beeinträchtigt ist, beziehungsweise falsche Inhalte vermittelt werden oder große inhaltliche Auslassungen gemacht werden.

Für eine bessere Übersicht werden die Fehler der einzelnen Dolmetscher:innen zunächst getrennt voneinander aufgelistet. Anschließend wird aus diesen einzelnen Fehleranalysen ein arithmetisches Mittel errechnet, der abschließend in Relation zu jenem Fehlerwert der maschinellen Dolmetschung gesetzt wird, um einen besseren insgesamten Vergleich zwischen Mensch und Maschine zu gewährleisten.

5 Ergebnisse

In diesem Kapitel werden die Ergebnisse der Analyse der maschinellen Dolmetschung durch die App Google Übersetzer, sowie der Dolmetschungen der Humandolmetscher:innen präsentiert. Anschließend sollen diese Ergebnisse zueinander in Beziehung gesetzt werden, um die eingangs gestellte Frage nach den Möglichkeiten und Grenzen von maschineller Dolmetschung mit Hilfe der App Google Übersetzer beantworten zu können. Die Aufarbeitung der Analyse der maschinellen Dolmetschung soll dabei zunächst anhand der einzelnen Schritte – Spracheingabe, maschinelle Übersetzung und Sprachausgabe –, die in Kapitel 3 genauer erläutert wurden, erfolgen, ehe die gesamte Dolmetschleistung des Programmes bewertet wird, um besser nachvollziehen zu können, bei welchen der Schritte es zu potentiellen Problemen kommen kann. Anschließend erfolgt die Bewertung der gesamten maschinellen Dolmetschung, wobei, wie im vorangegangenen Kapitel erläutert, die Übertragung des Inhalts und die Realisierung der Dolmetschung zunächst getrennt betrachtet werden, ehe eine Gesamtbewertung der maschinellen Dolmetschung erfolgt. In weiterer Folge werden die menschlichen Dolmetschleistungen nach demselben Prinzip bewertet, sodass als Abschluss dieses Kapitels die Gegenüberstellung der maschinellen und menschlichen Leistungen erfolgt und die Möglichkeiten und Grenzen von Google Übersetzer bei einer solchen maschinellen Konsekutivdolmetschung aufgezeigt werden können.

5.1 Spracheingabe

In diesem Abschnitt erfolgt die Analyse der Spracheingabe anhand des daraus von Google Übersetzer erstellten ausgangssprachlichen Transkriptes (T2), um daraus in weiterer Folge im Rahmen der Bewertung der gesamten Dolmetschung von Google Übersetzer ableiten zu können, an welchen Stellen im Prozess der maschinellen Dolmetschung welche konkreten Fehler zu beobachten sind.

5.1.1 Rechtschreibung und Zeichensetzung

Insgesamt fällt bei der Betrachtung des Transkriptes zunächst auf, dass die Groß- und Kleinschreibung nicht durchgehend korrekt erfolgt. Konkret gesagt, erfolgt am Satzanfang teilweise keine Großschreibung, wie beispielsweise direkt zu Beginn des Textes in den Zeilen 1 bis 3 (T2) erkennbar ist:

bonjour, je vous parle de la technologie de la reconnaissance faciale et surtout de son utilisation répandue en Chine et je commence. vous avez certainement déjà tous entendu parler de la technologie de la reconnaissance faciale, il s'agit d'une technologie qui permet de [...].^{4, 5}

An anderen Stellen wiederum, wie beispielsweise in Zeile 9 (T2): „Prenons l'exemple [...]“. [Nehmen wir das Beispiel [...]]. wird das erste Wort im Satz großgeschrieben, weshalb sich hier kein Muster erkennen lässt, unter welchen Bedingungen Großschreibung Anwendung findet, da diese sowohl am Anfang der automatisch erstellten Absätze als auch innerhalb dieser nur manchmal auftritt und keine Regelmäßigkeit zu erkennen ist.

Ein weiterer auffälliger Punkt ist die Setzung der Satzzeichen. Hierbei fehlen nicht nur einige Satzzeichen völlig, sondern es sind auch einige Satzzeichen an falschen Stellen gesetzt beziehungsweise es wurden nicht die jeweils richtigen Satzzeichen erkannt und verschriftlicht. Ein Beispiel für ein ausgelassenes Satzzeichen findet sich in Zeile 5 (T2), in welcher ein Fragesatz zu finden ist, der nicht nur anhand der Intonation im gesprochenen Ausgangstext, sondern auch anhand des Satzbaus beziehungsweise anhand des vorhandenen fakultativen Fragepartikels im Französischen eindeutig als solcher zu erkennen ist. So beginnt der Satz – wie auch Google Übersetzer korrekt transkribiert – mit „et pourquoi“ (T2: Z5) [und warum], also einem Fragewort, auf das der Fragepartikel „est-ce qu“ (T2: Z5) folgt. Allerdings wird am Ende dieses Satzes kein Fragezeichen gesetzt und auch auf keine andere Art eine Satzgrenze gezogen. Dadurch wird nicht nur dieser Fragesatz im Transkript nicht eindeutig als solcher erkennbar, sondern er verschmilzt aufgrund des gänzlich fehlenden abschließenden Satzzeichens zusätzlich mit dem folgenden Antwortsatz zu einem einzigen Satz, was zu Problemen bei der anschließenden maschinellen Übersetzung führen könnte. Vergleichbare Problematiken finden

4 guten Tag, je spreche mit Ihnen über die Technologie der Gesichtserkennung und vor allem von ihrer ausgebreiteten Verwendung in China und ich beginne. sie haben sicherlich schon von der Technologie der Gesichtserkennung gehört, es geht um eine Technologie, die uns erlaubt [...]

5 Bei den auf diese Art direkt angeführten Übersetzungen handelt es sich um möglichst wörtliche Übersetzungen der französischen Textstellen, die ausschließlich dazu dienen, das inhaltliche Verständnis der Analyse zu erhöhen, falls keine Französischkenntnisse bei den Leser:innen vorliegen sollten.

sich außerdem in den Zeilen 40, 60 und 71 (T2), in denen ebenso durch die Satzstellung beziehungsweise durch die Fragepartikel im Ausgangstext zu erkennen ist, dass es sich um Fragen handeln muss. So heißt es beispielsweise:

[...] mais ce qui est nouveau ici, c'est justement le recours à la reconnaissance faciale alors pourquoi est-ce que les citoyens chinois sont particulièrement inquiets de cette mesure et bien surtout parce qu'il craignent que leurs données ainsi enregistrées puisse être transmises [...]. (T2: Z38-41).⁶

Hier fehlt folglich einerseits der Punkt als Satzende vor dem Beginn des Fragesatzes mit „alors pourquoi est-ce que“ (T2: Z39-40), sowie andererseits auch das Fragezeichen am Ende dessen, sodass dieser Fragesatz auch mit dem folgenden Aussagesatz verbunden wird und somit in weiterer Folge vermutlich nicht richtig übersetzt werden kann.

Neben den auffallenden fehlenden Fragezeichen lässt sich außerdem feststellen, dass es auch bei den übrigen Satzgrenzen des Öfteren zu Problemen kommt. So werden im Transkript von Google Übersetzer einige Sätze zusammengefügt, wodurch teilweise sehr lange komplizierte und unübersichtliche Sätze entstehen, wie beispielsweise in Zeile 9 bis 12 (T2) zu sehen ist:

Prenons l'exemple par exemple des téléphones portables que l'on peut déverrouiller grâce à la reconnaissance faciale cela permet à l'utilisateur de ne pas devoir reconnu de ne pas devoir se souvenir d'un mot de passe compliqué et de ne pas avoir à l'entrée l'encoder chaque fois qu'il veut utiliser son téléphone.⁷

Auch hier kann anhand des Satzbaus erkannt werden, dass es sich um mehrere Sätze handeln muss, die durch einen Punkt getrennt werden müssten. So müsste der erste Satz mit „que l'on peut déverrouiller grâce à la reconnaissance faciale“ (T2: Z9-10) enden, da danach „cela permet à l'utilisateur de ne pas devoir [...]“ (T2: Z10) folgt, was syntaktisch, grammatikalisch und inhaltlich eindeutig einen Satzanfang markiert.

An anderer Stelle (T2: Z68) wiederum finden sich falsch gesetzte Punkte, sodass ein einziger Satz in drei Sätze aufgeteilt wurde, wodurch es abermals zu Problemen bei der maschinellen Übersetzung kommen könnte, da dadurch falsche textuelle Verknüpfungen entstehen, wodurch die Kohärenz beeinträchtigt werden kann beziehungsweise falsche inhaltliche Zusammenhänge entstehen könnten. So wurde hier etwa die Nominalphrase „des directives.

⁶ [...] aber was hier neu ist, ist einfach das Zurückgreifen auf die Gesichtserkennung also warum sind die chinesischen Bürger besonders besorgt über diese Maßnahme und gut vor allem, weil sie befürchten, dass ihre so gespeicherten Daten übermittelt werden könnten [...].

⁷ Nehmen wir zum Beispiel Mobiltelefone, die man dank der Gesichtserkennung entsperren kann das erlaubt den Nutzern sich keine komplizierten Passwörter merken zu müssen und keinen Eingang Code jedes Mal zu haben, wenn man das Telefon benutzen möchte.

Non contraignants“ (T2: Z68) [Regeln. Nicht verbindlich] fälschlicherweise durch einen Punkt getrennt. Außerdem erfolgt auch davor durch eine fehlerhafte Satzzeichensetzung eine falsche Verknüpfung von zwei Sätzen. So heißt es an jener Textstelle wie folgt:

En Chine on compte environ 850 millions d'utilisateurs de téléphones portable, ça constitue une précieuse base de données pour les entreprises et les administrations alors pour tenter de rassurer la population. Le gouvernement a publié des directives. Non contraignantes visant réglementer la collecte et l'utilisation des données.⁸ (T2 : Z66-69).

Hierbei endet jedoch der erste Satz nach „les administrations[.]“ (T2: Z67) und nicht nach „la population.“ (T2: Z68). Durch diese Satzzeichensetzung wird die Infinitivkonstruktion „pour tenter [...]“ (T: Z67) fälschlicherweise mit dem Satz zuvor anstatt mit dem Folgenden verknüpft, wodurch der Inhalt im Transkript falsch wiedergegeben wird.

Auch die Beistrichsetzung seitens Google Übersetzer ist fehlerbehaftet. Am auffallendsten dabei ist jener Beistrich in der Zeile 63 (T2), der zwischen dem Subjekt und dem anschließenden Prädikat inklusive Negationspartikel gesetzt wurde, was auch bei der Übersetzung zu Problemen führen könnte, wenn das Programm das Subjekt und das Prädikat dadurch nicht richtig miteinander in Verbindung bringen kann. Außerdem fehlen an vielen Stellen Beistriche, wie etwa an der eben genannten Stelle, wo ein obligatorischer Beistrich nach dem vorangestellten Konnektor „par conséquent“ (T2: Z62) [Folglich] fehlt. Ähnliches findet sich in einigen weiteren Textstellen, wobei die Beistrichsetzung andererseits teilweise auch korrekt durchgeführt wurde, weshalb sich auch bei der Zeichensetzung insgesamt – ähnlich wie bei der Groß- und Kleinschreibung – nicht wirklich erkennen lässt, nach welchen Regeln diese vollzogen wird und warum diese an manchen Stellen problemlos funktioniert, während sie an anderen Stellen recht fehlerhaft erscheint, während die Sprecherin der Ausgangsrede hinsichtlich Satzmelodie und Prosodie über die gesamte Rede hinweg gleichbleibend spricht, sodass sich hier kein eindeutiger Grund feststellen lässt.

⁸ In China zählt man ungefähr 850 Millionen Mobiltelefonnutzer, das bildet eine wertvolle Basis an Daten für die Unternehmen und die Verwaltungen also, um zu versuchen, die Bevölkerung zu beruhigen. Die Regierung hat Richtlinien veröffentlicht. Nicht verbindlich, die das Sammeln und das Verwenden von Daten reglementieren sollen.

5.1.2 Verzögerungslaute und Selbstkorrekturen

Einen weiteren Punkt, der wie in Kapitel 3 beschrieben zu Problemen bei der Spracherkennung führen kann, stellen Verzögerungslaute dar, da diese zum Beispiel fälschlicherweise als semantische Lexeme transkribiert werden könnten oder aber auch die Erkennung von Wörtern in deren Umgebung erschweren. Im Transkript lässt sich jedoch erkennen, dass Verzögerungslaute kein allzu großes Problem hinsichtlich der Transkripterstellung darstellen, was unter anderem wohl auch daran liegt, dass diese im Ausgangstext bereits nur sehr selten vorkommen. Tatsächlich finden sich im Transkript lediglich zwei Stellen, an denen ein Verzögerungslaut fälschlicherweise als Lexem transkribiert wurde. So steht in der Zeile 11 (T2) „à“ [an], sowie in der Zeile 40 (T2) „et“ [und], wobei im Ausgangstext an diesen beiden Stellen lediglich ein nicht lexikalischer Verzögerungslaut artikuliert wird, was folglich bei der Übersetzung wiederum zu Schwierigkeiten führen könnte.

Neben den Verzögerungslauten können auch Selbstkorrekturen ein Problem für die maschinelle Dolmetschung darstellen, da diese einerseits bei der Spracherkennung bereits zu Problemen führen können, wenn es sich beispielsweise um nicht vollständig artikuliert Wörter handelt oder aber auch wenn im zu übersetzenden Text Doppelungen vorkommen, die für die maschinelle Übersetzung eine Herausforderung darstellen können. Im vorliegenden Transkript finden sich zwei Doppelungen, die auf Selbstkorrekturen im Ausgangstext zurückzuführen sind. Diese finden sich in den Zeilen 10-11 und 79 (T2), sodass diese Stellen im Transkript wie folgt lauten:

[...] cela permet à l'utilisateur de ne pas devoir reconnu de ne pas devoir se souvenir d'un mot de passe compliqué et de ne pas avoir à l'entrée l'encoder chaque fois qu'il veut utiliser son téléphone.⁹ (T2: Z10-11)

sowie „les citoyens. les citoyens chinois“ (T2 : Z79) [die Bürger. Die chinesischen Bürger]. Außerdem wurde etwa in Zeile 18 (T2) „favorable à recours“¹⁰ transkribiert, wobei sich im Ausgangstext an dieser Stelle eine Selbstkorrektur wiederfindet, die wie folgt lautet: „favorables à/ au recours“ (T1: Z20)¹¹. Hierbei wurde also die Selbstkorrektur seitens der Spracherkennung nicht transkribiert, weswegen sich hier nur die zuerst artikuliert Version wiederfindet. Da es sich bei dieser Textstelle ohne die Selbstkorrektur aber lediglich um einen leichten

⁹ [...] dies erlaubt den Nutzern, sich kein kompliziertes Passwort merken zu müssen und keinen Eingang Code jedes Mal zu haben, wenn man sein Telefon benutzen möchte.

¹⁰ positiv eingestellt auf Zurückgreifen

¹¹ positiv eingestellt auf das Zurückgreifen

Grammatikfehler handelt, sollte die Auswirkung davon auf die maschinelle Übersetzung nicht allzu schwerwiegend sein.

5.1.3 Homophonie

Eine weitere Schwierigkeit bei der Spracherkennung stellen Homophone dar, die gerade im Französischen recht häufig vorkommen, da viele Phoneme grundsätzlich auf mehrere Arten in der Orthographie wiedergegeben werden können (vgl. Champoux 2018: 65). Die falsche Transkription von Homophonen kann, wie in Kapitel 3 erläutert wurde, zu grammatikalischen Ungenauigkeiten beziehungsweise Fehlern oder aber auch zur Transkription von anderen Lexemen führen, wodurch die Bedeutung einer Aussage verändert werden kann. Im Transkript der Spracheingabefunktion von Google Übersetzer finden sich einige Transkriptionsfehler, die auf Homophonie zurückzuführen sind und mehr oder weniger starke Sinnverzerrungen des Textes bewirken. So finden sich etwa in den Zeilen 16, 18 und 70 (T2) Transkriptionsfehler, die auf grammatikalische Homophonie zurückzuführen sind und die Übereinstimmung zwischen Adjektiven beziehungsweise Partizipien und den dazugehörigen Substantiven betreffen.

So heißt es beispielsweise „La reconnaissance faciale est en effet déjà confortablement installé dans le quotidien des citoyens chinois“ (T2: Z16-17)¹², wobei hier das Partizip ‚installé‘ nicht an das feminine Genus des Substantives ‘reconnaissance’ angepasst wurde, sodass es eigentlich ‚installée‘ lauten müsste. Weiters findet sich die Phrase „les Chinois [...] semblent assez favorable“¹³ (T2: Z18), in der das Adjektiv fälschlicherweise nicht in der Pluralform mit dem im Mündlichen stummen Pluralmorphem <s> transkribiert wurde. In der Zeile 70 (T2) heißt es schließlich „de vrais législations“, sodass auch hier das Adjektiv hinsichtlich des femininen Genus des Substantives nicht richtig dekliniert wurde. Bei all diesen Textstellen des Transkriptes von Google Übersetzer handelt es sich folglich zwar um Grammatikfehler, die aber für die maschinelle Übersetzung kein allzu großes Problem darstellen sollten, sofern die inhaltliche Zusammengehörigkeit der jeweiligen Wörter beziehungsweise Wortgruppen trotzdem erkannt werden kann.

Zusätzlich finden sich einige Probleme bei der korrekten Transkription von Pluralformen, da es sich hierbei im Französischen auch oft um Homophone handelt. Beispiele hierfür

¹²) Die Gesichtserkennung ist eigentlich schon komfortabel im Alltag der chinesischen Bürger eingerichtet

¹³ die Chinesen [...] scheinen so positiv

finden sich etwa in den Zeilen 32 bis 34 (T2), wo aus dem Kontext heraus klar ist, dass es sich bei „leur visage“ (T2: Z32) [ihr Gesicht], „leur numéro“ (T2: Z33) [ihre Nummer] und „leur visage“ (T2: Z33-34) [ihr Gesicht] eigentlich um Pluralformen handelt, an die demnach jeweils das stumme Pluralmorphem <s> suffigiert werden müsste.

Ähnliches findet sich in der Textstelle „il craignent que leurs données ainsi enregistrées puisse être transmises et revendus“¹⁴ in Zeile 41 (T2). Hier wurde zunächst das homophone Pronomen nicht an die Pluralform des Verbs angepasst, die sich wiederum lautlich von der Singularform unterscheidet, weswegen diese wohl richtig transkribiert wurde. Außerdem wurde das Adjektiv weder an die Pluralform noch an das feminine Genus des Substantivs angepasst. Auch eines der Partizipien wurde aufgrund von Homophonie nicht an das feminine Genus angepasst. Zusätzlich findet sich auch das Verb in der Singularform wieder, obwohl es sich beim Subjekt um eine Pluralform handelt. Ein weiteres Beispiel steht in Zeile 46 (T2), in der das Substantiv „numéro“ [Nummer] nicht an den vorangegangenen Artikel im Plural angepasst wurde, obwohl durch diesen der Plural auch lautlich eigentlich eindeutig markiert wurde.

Neben diesen Formen der grammatikalischen Homophonie finden sich im Transkript allerdings auch einige lexikalische Homophone, deren fehlerhafte Transkription zu Sinnverfälschungen bei der maschinellen Übersetzung führen kann. Ein Beispiel hierfür findet sich in Zeile 25 (T2), in welcher „des citoyens des obéissants“ [die Bürger die Gehorsamen] transkribiert wird, wobei im Ausgangstext „des citoyens désobéissants“ (T1: Z27) [die ungehorsamen Bürger] gesagt wird. Hierbei handelt es sich zwar um zwei homophone Ausdrücke, jedoch ist die Transkription von Google Übersetzer ungrammatisch, weswegen bei Betrachtung des grammatikalischen Kontextes klar wird, um welche Form es sich eigentlich handeln muss. Der Fehler hier ist besonders schwerwiegend, da es sich um Antonyme – gehorsam und ungehorsam – handelt, weswegen es im Zuge der maschinellen Dolmetschung an dieser Stelle wohl zu groben Problemen kommen wird, da eine Sinnumkehr aufgrund der fehlerhaften Spracherkennung wahrscheinlich erscheint.

Ein weiteres Beispiel für eine ähnliche fehlerhafte Transkription findet sich in Zeile 46 und 48 (T2), wo statt „données“ (T1: Z50, Z52) [Daten] „donner“ (T2: Z46, Z48) [geben] transkribiert wurde. Hierbei gilt jedoch ähnliches wie bei dem oben erwähnten Beispiel, sodass aus dem syntaktischen und grammatikalischen Kontext heraus eigentlich klar ersichtlich ist, um welche Form der Homophone es sich handeln muss, da der so transkribierte Satz ungrammatisch ist. Auffallend ist dabei auch, dass es in den Zeilen 41 und 43 (T2) nicht zu diesem

¹⁴ sie befürchten, dass ihre so gespeicherten Daten übermittelt und verkauft werden könnten

Problem gekommen ist und „données“ korrekt wiedergegeben wurde, obwohl in allen erwähnten Fällen der syntaktische Kontext ähnlich ist. Demnach ist abermals nicht klar ersichtlich, warum es in manchen Fällen zu derartigen Problemen kommt, während die Spracherkennung in Bezug auf solche Homophone an anderen Stellen gut zu funktionieren scheint. In diesem Fall handelt es sich bei den Homophonen zwar nicht um Antonyme, aber trotzdem um Lexeme mit anderem semantischem Inhalt, weswegen es auch hier zu Missverständnissen bei der maschinellen Übersetzung kommen könnte.

Weitere Fehler bezüglich Homophonie finden sich in Zeile 26 (T2), in welcher „voir“ [sehen] statt „voire“ (T1: Z29) [sogar], sowie in Zeile 55 (T2), in welcher statt „a décidé“ (T1: Z60) [hat entschieden], „à décider“ [zu entscheiden] transkribiert wurde. Hier gilt Ähnliches wie oben, sodass die Wahl der richtigen Schriftform und somit der richtigen Bedeutung des Wortes aus dem syntaktischen Kontext erschlossen werden kann und es durch die unterschiedliche semantische Bedeutung der beiden Lexeme beziehungsweise die unterschiedliche grammatikalische Funktion wiederum zu Problemen bei der Übertragung des Inhalts im Rahmen der maschinellen Dolmetschung kommen könnte.

Zusätzlich kommen einige weitere Fehler vor, die zwar auf keine vollständige Homonymie, aber auf recht ähnlich klingende Wörter beziehungsweise Phrasen zurückzuführen sind. Ein Beispiel hierfür findet sich in Zeile 42 (T2), wo „par“ [durch] transkribiert wurde, wobei es im Ausgangstext „pas“ (T1: Z46) [nicht] lautet. Dies ist insofern von großer Relevanz, da es sich im Ausgangstext um einen Negationspartikel handelt, welcher von der Spracherkennung als anderes Lexem – in diesem Fall als eine Präposition – erkannt wurde, wodurch sich wohl eine starke Sinnverfälschung bei der maschinellen Übersetzung ergeben wird. Auch hier gilt wiederum, dass sich aus dem syntaktischen Kontext eindeutig erschließen lässt, dass es sich nicht um das von Google Übersetzer transkribierte Wort handeln kann, da es sich sonst um einen ungrammatischen Satz handeln würde.

Weitere Fehler hinsichtlich ähnlich klingender Wörter und Phrasen, die allerdings keinen so starken Bedeutungsunterschied wie das eben erläuterte Beispiel darstellen, finden sich etwa in Zeile 35 (T2), in der das Pronomen „l“ [es] eingefügt wurde, das im Ausgangstext nicht vorhanden ist und zu dem auch das Bezugswort fehlt. In Zeile 29 (T2) hingegen wurde der Artikel, in Zeile 42 (T2) das Verb „est“ [ist] ausgelassen und in Zeile 49 (T2) „un“ [ein] statt „a en“ (T1: Z54) [hat davon] und in Zeile 71 „dire“ [sagen] statt „dira (T1: Z79) [wird sagen] transkribiert. Bei diesen Fehlern handelt es sich abermals lediglich um leichte Grammatikfehler, bei denen es abzuwarten bleibt, ob diese einen Einfluss auf die Qualität der maschinellen Übersetzung haben werden. Zusätzlich findet sich „à l’heure“ (T2: Z69) [rechtzeitig]

anstelle von „alors“ (T1: Z77) [also], was eine andere lexikalische Bedeutung hat, für die Übersetzung aber von nicht allzu großer Relevanz sein sollte.

Weiters findet sich die Transkription „viaen“ (T2: Z71) statt „bien“ (T1: Z79) [gut], was interessant zu betrachten ist, da es sich bei dem hier transkribierten Wort um kein existierendes französisches Wort handelt, während das hier artikulierte Wort im Ausgangstext ein sehr gewöhnliches Füllwort darstellt. Demnach trägt das Wort im Ausgangstext auch keine relevante Bedeutung, die eventuell durch eine falsche Übersetzung verloren gehen könnte. Dennoch gilt es zu überprüfen, inwiefern diese Wortneuschöpfung der Spracherkennung durch die maschinelle Übersetzung übertragen wird.

5.1.4 Zusammenfassung

Zusammenfassend lässt sich zur Analyse der Spracherkennungsfunktion von Google Übersetzer feststellen, dass es sich insgesamt um ein recht gutes Transkript handelt, das den Inhalt über weite Teile korrekt wiedergeben kann. Allerdings finden sich einige grammatikalische Fehler, die auf die häufigen Homophone im Französischen zurückzuführen sind und darauf schließen lassen, dass bei der Spracherkennung der syntaktische Kontext einzelner Wörter nicht ausreichend miteinbezogen wird, da sich die auftretenden Fehler dadurch vermeiden lassen würden. Hinzu kommt, dass viele Satzzeichen nicht richtig gesetzt sind, sodass diese entweder gänzlich fehlen, an den falschen Stellen vorkommen oder die falsche Art von Satzzeichen gewählt wurde, sodass beispielsweise anstatt eines Fragezeichens ein Punkt aufscheint. Auch hier würden sich die meisten Fehler durch stärkere Einbeziehung des syntaktischen Kontextes oder durch eine genauere Analyse der Intonation und Prosodie des Ausgangstextes wohl vermeiden lassen. Demnach könnte das Transkript in dieser Rohfassung wohl keine professionelle Transkription ersetzen und es ist demnach auch zu vermuten, dass sich die Fehler in dem Transkript auf die Qualität der maschinellen Übersetzung auswirken können, worauf im folgenden Kapitel genauer eingegangen wird.

Außerdem ist anzumerken, dass es sich bei dem in mündlicher Form vorliegenden Ausgangstext um einen Text in einem idealisierten Setting handelt, wo keine technischen Probleme und kaum außersprachliche Störgeräusche vorliegen, die eine Spracherkennung zusätzlich erschweren würden. Hinzu kommt, dass wenig Verzögerungslaute und Selbstkorrekturen seitens der Sprecherin vorkommen, sodass es auch hier zu wenig Problemen kommt und die

Spracherkennung auch nahezu alle auftretenden Verzögerungslaute herausfiltern konnte. Lediglich die Selbstkorrekturen wurden so im Transkript übernommen, was abermals zu Beeinträchtigungen bei der maschinellen Übersetzung führen kann. Auch der nonverbale Input der Rednerin spielt in diesem Fall keine tragende Rolle und der Text enthält keine Fremdwörter, Neologismen oder Eigennamen, die eine zusätzliche Schwierigkeit für die Spracherkennung darstellen würden und Gegenstand für zukünftige Forschung sein könnten. Somit bleiben als größtes Hindernis die korrekte Erkennung von Homonymen sowie das Setzen von Satzzeichen.

5.2 Maschinelle Übersetzung

In diesem Kapitel erfolgt die Analyse der maschinellen Übersetzung des zuvor erstellten Transkriptes durch Google Übersetzer, ehe im folgenden Kapitel der abschließende Schritt der Sprachausgabe erarbeitet wird.

5.2.1 Situatives Wissen und innertextliche Verweise

Insgesamt lässt sich bei der Betrachtung der maschinellen Übersetzung festhalten, dass dem Programm – wie in Kapitel 3 bereits erläutert wurde – das situative Wissen fehlt. Dementsprechend ergibt sich etwa das Problem, dass im Zieltext nicht gegendert wird, was im Deutschen üblicher wäre und auch von den menschlichen Dolmetscher:innen teilweise so gemacht wurde, wie in Kapitel 6.5 genauer erläutert wird. In der französischsprachigen Ausgangskultur hingegen ist Gendern noch kein Usus, weswegen es seitens der Dolmetschenden durchaus angebracht wäre, dies für die Zielkultur anzupassen. Hinzu kommt, dass im Ausgangstext die Rezipient:innen gesiezt werden, wobei dieses Pronomen im Französischen auch jenem für die Anrede der zweiten Person Plural entspricht, wodurch sich eine gewisse Doppeldeutigkeit ergibt, die jedoch von der maschinellen Übersetzung richtig übertragen wurde, sodass im deutschen Zieltext auch die passende Höflichkeitsform steht. Eine Ausnahme davon stellt wohl die Anrede zu Beginn des Ausgangstextes „bonjour“ (T2: Z1) [guten Tag] dar, die im Zieltext mit „Hallo“ (T3: Z1) übersetzt wurde, wobei dies im Deutschen etwas informeller als im französischsprachigen Ausgangstext wirkt.

Ein weiteres allgemeines Problem, das sich bei maschineller Übersetzung ergeben kann, sind fehlerhaft hergestellte transphrastische Bezüge, also innertextliche Verweise. Hierbei findet sich in den Zeilen 16, 24 und 37 (T3) der fehlerhafte Verweis „es“, der in Bezug zu „Gesichtserkennung“ steht und dementsprechend in der femininen Form vorliegen müsste. Dabei ist interessant, dass es sich im französischsprachigen Ausgangstext zufälligerweise ebenfalls um ein feminines Substantiv handelt, auf das dementsprechend mit einem femininen Pronomen verwiesen wird, weshalb hier eher nicht davon auszugehen ist, dass diese Referenz wörtlich übernommen wird, was in diesem Fall sogar korrekt wäre. Stattdessen wird im Deutschen die Referenz mit einem neutralen Pronomen hergestellt, weshalb es sich um einen grammatikalischen Fehler handelt.

5.2.2 Lexikalische und syntaktische Ambiguitäten

Einen weiteren Punkt, der zu Schwierigkeiten bei der maschinellen Übersetzung führen kann, stellen lexikalische oder syntaktische Ambiguitäten dar. Hierbei findet sich im Text das Beispiel des Wortes „encoder“ (T2: Z11), das mit „verschlüsseln“ (T3: Z10) übersetzt wurde. Es steht hier im Kontext mit dem Passwort des Handys, weswegen das angesprochene Passwort nicht verschlüsselt werden muss, sondern es vielmehr darum geht, das Passwort nicht mehr benutzen zu müssen und sein Handy ohne dieses entsperren zu können. Dies wird auch bei den Textproduktionen der Humandolmetscher:innen sichtbar, worauf in Kapitel 6.5 noch genauer eingegangen wird.

Hinsichtlich der Unterschiede beim Satzbau, die zu fehlerhaften Übersetzungen führen, finden sich keine wirklichen direkt darauf zurückzuführenden Probleme im Text. Vielmehr lassen sich jene Sätze, die im Zieltext Fehler hinsichtlich korrekter Satzstellung aufweisen, auf andere Problematiken der Spracherkennung zurückführen, wie etwa falsche Zeichensetzung oder die falsche Erkennung von Homonymen, wie im vorangegangenen Kapitel bereits erwähnt wurde. Aus diesem Grund soll auf die weiteren Fehler, die sich im inhaltlichen Vergleich mit dem Ausgangstext erkennen lassen, aber lediglich auf dem fehlerhaften Transkript der Spracherkennung beruhen, in Kapitel 6.4 im Rahmen der Bewertung der gesamten Dolmetschung genauer eingegangen werden, damit im weiteren Verlauf dieses Kapitel vor allem auf jene Problematiken eingegangen werden kann, die rein der Funktionsweise der maschinellen Übersetzung geschuldet sind.

5.2.3 Einfügungen und Auslassungen

Ein Fehler, dessen Ursprung zwar in Verbindung mit einem an der falschen Stelle gesetzten Satzzeichen durch die Spracherkennung steht, dessen Wirkung aber durch die maschinelle Übersetzung verändert wurde, findet sich in den Zeilen 39 bis 42 (T3):

Warum also sind die chinesischen Bürger besonders besorgt darüber? Maßnahme und vor allem weil sie befürchten, dass ihre so erfassten Daten tatsächlich an andere Unternehmen übermittelt und weiterverkauft werden könnten, ist es in China tatsächlich selten, dass [...]. (T3: Z39-42)

Hierbei wurde im Zuge der Übersetzung zwar das Fragezeichen und das syntaktisch nötige Adverb „darüber“ (T3: Z40) eingefügt, welche im französischsprachigen Transkript noch nicht vorhanden waren, allerdings steht das Fragezeichen in der Übersetzung an der falschen Stelle. Deswegen wurde zwar der Fragesatz grammatikalisch und inhaltlich korrekt übersetzt, der darauffolgende Satz ist allerdings inhaltlich und grammatikalisch unstimmig, da das Wort „Maßnahme“ (T3: Z40) keinen direkten Bezug zu diesem hat, da diese eigentlich noch innerhalb des Fragesatzes stehen müsste.

Zudem finden sich einige Fehler, die darauf basieren, dass während der Übersetzung einzelne Worte oder Phrasen nicht übersetzt, sondern stattdessen ausgelassen wurden. Dies ist beispielsweise in Zeile 2 (T2) der Fall, wo „déjà“ [schon] zwar transkribiert wurde, in der Übersetzung allerdings nicht aufscheint. Ähnliches gilt für „en plus“ (T2: Z20) [zusätzlich], wobei bei diesen beiden Beispielen der Inhalt des Ausgangstextes trotz der Auslassungen richtig übertragen werden kann, da es sich lediglich um Satzadverbien handelt, die in diesen Fällen keine allzu starke lexikalische Bedeutung in sich tragen.

Etwas anders sieht es bei der Auslassung von „sujet“ (T2: Z70) [Thema] aus, wobei hier nicht ersichtlich wird, warum dies in der Übersetzung weggelassen wurde. So handelt es sich im ausgangssprachlichen Transkript von Google Übersetzer (T2: Z69-70) um einen vollständig transkribierten, grammatikalisch und syntaktisch nahezu korrekten und logisch schlüssigen Satz, der auch mit einem Punkt am Ende abgeschlossen wird: „Mais à l’heure actuelle la Chine ne dispose pas de vrais législations à ce sujet.“ (T2: Z69-70)¹⁵ In der Übersetzung jedoch wird nahezu der gesamte Satz richtig übertragen, allerdings endet der Satz unvollständig mit „[...] keine wirkliche Gesetzgebung zu diesem“ (T3: Z69-70) und ohne Punkt, sodass die

¹⁵ Aber aktuell hat China keine wirklichen Gesetzgebungen zu diesem Thema.

Übersetzung von „sujet.“ (T2: Z70) fehlt, was sich rein durch Betrachtung des Transkriptes nicht erklären lässt.

Eine weitere Auslassung findet sich bei „à décider“ (T2: Z55) [zu entscheiden]. Hierbei kommt in der Transkription ein grammatikalischer Fehler vor, da es sich eigentlich um eine finite und keine infinite Verbform handeln müsste. Dieser Fehler lässt sich mit Homonymie begründen, wie im vorangegangenen Kapitel erläutert wurde. Das Fehlen des finiten Verbes in diesem Satz könnte der Grund für die Auslassung sein, wobei anzumerken ist, dass dieser Satz (T3: Z54-56) durch die maschinelle Übersetzung zu einem grammatikalisch korrekten Satz in der Zielsprache umformuliert wurde und stattdessen das weitere infinite Verb des Ausgangstextes „refuser“ (T2: Z55) [verweigern] zu einem finiten Verb geändert wurde. Dadurch wurde zwar die inhaltliche Aussage des Satzes leicht verändert, dafür allerdings ein ungrammatischer Satz des Ausgangstextes durch die maschinelle Übersetzung in Bezug auf die Grammatik richtiggestellt.

5.2.4 Umformulierungen

Zusätzlich finden sich einige unpassende Formulierungen im Zieltext, die nicht auf Fehler in der Spracherkennung zurückzuführen sind, sondern rein auf der Funktionsweise der maschinellen Übersetzung beruhen, da diese im von Google Übersetzer angefertigten französischsprachigen Transkript noch korrekt wiedergegeben wurden. Dazu zählt beispielsweise die Übersetzung von „on est forcé“ (T2: Z13), die „wir sind gezwungen“ (T3: Z11-12) lautet. Dies erscheint in diesem Kontext im Deutschen unpassend, da es andere Konnotationen weckt. Die Bedeutung im Ausgangstext ist eher im übertragenen Sinn zu verstehen, wie sich auch an den Dolmetschungen der menschlichen Translator:innen erkennen lässt und hier etwa mit „[es] bleibt einem gar nichts anderes übrig“ (T8: Z13) oder „dann muss man“ (T7: Z13) gedolmetscht wurde.

Ähnliches gilt auch für „confortablement installé“ (T2: Z16), was sich in diesem Fall auf die Gesichtserkennungstechnologie bezieht, weswegen die recht wörtliche deutsche Übersetzung „bequem [...] installiert“ (T3: Z15-16) ebenso unpassend erscheint, da auch ‚installieren‘ im Deutschen eine andere Konnotation als im Französischen aufweist. Die Übersetzung müsste hier stärker vom Ausgangstext abweichen, um einen adäquaten Zieltext produzieren zu können. Selbiges gilt für die Übersetzung von „synonyme“ (T2: Z20) mit „gleichbedeutend“ (T3: Z20), die ebenfalls in dem hier auftretenden Kontext nicht wirklich passend ist. Außerdem

findet sich die Übersetzung „Führer“ (T3: Z62) im Zieltext, wozu ebenfalls anzumerken ist, dass dieser Begriff im Deutschen eine deutlich andere Konnotation als der Begriff „leader“ (T1: Z62) im Ausgangstext aufweist und demnach im Deutschen eher vermieden werden sollte, wie auch bei den Humandolmetschungen sichtbar wird.

Weitere Formulierungen, die im Deutschen eher ungewöhnlich klingen, aber im Gegensatz zum eben genannten Beispiel keine so stark abweichende Konnotation aufweisen, finden sich in Zeile 50 (T3) mit „Mobiltelefonanwendungen“ statt etwa ‚Apps‘ oder in Zeile 53 (T3) mit der recht wörtlichen Übersetzung „in jüngster Zeit ist es ein weiteres Beispiel“.

Allerdings findet sich mit der Phrase „[sie müssen] ihr Gesicht registrieren, um erkannt zu werden zur Gesichtserkennung“ (T3: Z31-32) im Vergleich zu „[ils] doivent faire enregistrer leur visage pour pouvoir être reconnu grâce à la reconnaissance faciale“ (T2: Z31-32) [[sie] müssen ihr Gesicht registrieren lassen, um dank der Gesichtserkennung wiedererkannt werden zu können] auch ein Beispiel dafür, dass eine Übersetzung nicht wörtlich genug gemacht wurde, obwohl die von Google Übersetzer erstellte Transkription korrekt ist. Dadurch entsteht abermals ein ungrammatischer Satz im Deutschen, der demnach auch den ursprünglichen Sinn des Ausgangstextes im Zieltext nicht korrekt wiedergeben kann.

5.2.5 Syntax

In Bezug zu fehlerhaftem Satzbau in der maschinellen Übersetzung finden sich ebenso einige Beispiele. Dazu zählen etwa die Zeilen 58-61 (T3):

[...] aber was die chinesische Regierung dazu treibt, Gesichtserkennung in China durchzusetzen alle Aspekte des täglichen Also offiziell offensichtlich, wie ich sagte, dies ist eine Maßnahme, um die Sicherheit der Bürger zu gewährleisten, aber [...]. (T3: Z58-61)

Hierbei handelt es sich im Ausgangstext um einen Fragesatz und einen darauffolgenden Aussagesatz. Durch den oben bereits erläuterten Fehler beim Prozess der Spracherkennung ist zwar der Satzbau in beiden Sätzen so vorhanden, allerdings fehlt das Fragezeichen. In der Übersetzung jedoch passt die Wortstellung nicht für einen Fragesatz im Deutschen, wozu kommt, dass die Satzgrenze zum darauffolgenden Satz nicht mehr gegeben ist, da der Punkt am Satzende wegfällt, obwohl das eigentlich erste Wort des Folgesatzes mit einem Großbuchstaben geschrieben wird. Zusätzlich wurde das lexikalisch ambige Wort „quotidien“ (T2: Z60) [Alltag / alltäglich], welches prinzipiell ein Adjektiv oder ein Substantiv darstellen kann, seitens der

maschinellen Übersetzung an dieser Stelle als Adjektiv anstelle eines Substantivs interpretiert, weswegen dieses im Zieltext fehlt. Dadurch entsteht im Deutschen ein unvollständiger und tendenziell unverständlicher Satz.

Ein weiteres Problem hinsichtlich der Satzstellung findet sich in den Zeilen 61-64 (T3) des Zieltextes:

[...] aber Sie sollten auch wissen, dass die chinesische Regierung daran interessiert ist, China zu einem technologischen Führer zu machen, deshalb die Regierung, n zögern Sie nicht, Gesichtserkennung oder künstliche Intelligenz zu unterstützen Programme im Allgemeinen. (T3: Z61-64)

Hier entsteht ein ebenso unvollständiger, nahezu unverständlicher beziehungsweise logisch inkohärenter Satz durch die maschinelle Übersetzung, obwohl der Satz im Transkript von Google Übersetzer (T2: Z62-64) bis auf einen zusätzlichen Beistrich zwischen Subjekt und Prädikat richtig transkribiert wurde, weswegen anzunehmen ist, dass dieser vermeintliche Beistrich die Funktionsweise der maschinellen Übersetzung stark störte. So findet sich beispielsweise der französische Negationspartikel „n“ (T2/T3: Z63) direkt übertragen in der Übersetzung wieder, obwohl dies im deutschen Zieltext klarerweise nicht passend ist. Außerdem wurde das finite Verb, das im Ausgangstext in der Indikativform der dritten Person flektiert ist, zu einer Imperativform der Höflichkeitsform verändert und hat somit keinen Bezug zum eigentlichen Subjekt mehr. Demnach ist das ursprüngliche Subjekt „die Regierung“ (T3: Z63) in der Übersetzung zwar vorhanden, allerdings ohne wirkliche syntaktische Funktion, wodurch sich die fehlerhafte syntaktische Struktur des Satzes ergibt. Hinzu kommt, dass im zweiten Abschnitt des Satzes das Kompositum des Ausgangstextes „des programmes de reconnaissance faciale ou d’intelligence artificielle en général“¹⁶ (T2: Z63-64) im Zuge der maschinellen Übersetzung aufgelöst wurde und als solches im Zieltext nicht mehr vorhanden ist. So wurden die einzelnen Bestandteile dieses Kompositums einzeln übersetzt, wobei allerdings der Zusammenhang zwischen diesen und damit auch der eigentliche Inhalt verloren gehen. Dementsprechend ist der betreffende Teil (T3: Z63-64) des Zieltextes nicht nur ungrammatisch, sondern hat aufgrund der falsch hergestellten Bezüge der einzelnen Wörter auch nicht mehr die Wirkung, die er noch im Ausgangstext innehatte.

¹⁶ Programme zur Gesichtserkennung oder zur künstlichen Intelligenz im Allgemeinen

5.2.6 Verbesserungen

Andererseits finden sich in der Übersetzung auch Textpassagen, die im Vergleich zum von Google Übersetzer erstellten Transkript eine Verbesserung darstellen, was bedeutet, dass Fehler, die im französischen Transkript zu finden sind, im deutschen Zieltext in dieser Art nicht mehr vorhanden sind. Ein Beispiel hierfür findet sich in Zeile 27 (T3) der Übersetzung. Hier wurde seitens der Spracherkennung ein falsches Homonym mit einer anderen lexikalischen Bedeutung („voir“ (T2: Z26) [sehen] statt „voire“ (T1: Z29) [sogar]) transkribiert, welches grammatikalisch auch nicht an dieser Stelle in den Satz passt, wie im vorangegangenen Kapitel ausführlicher erläutert wurde. In der Übersetzung jedoch taucht wiederum die richtige Formulierung „sogar“ (T3: Z 27) anstelle von ‚sehen‘ auf, was die eigentliche wörtliche Übersetzung laut dem von Google Übersetzer erstellten Transkript wäre. An dieser Stelle konnte durch die maschinelle Übersetzung also der durch die Spracherkennung entstandene Fehler ausgebessert werden, sodass der ursprüngliche Sinn des Ausgangstextes transferiert werden konnte und wohl angenommen werden kann, dass für diese Übersetzung auch der Kontext miteinbezogen wurde, da anderenfalls eine solche Übersetzung nicht möglich wäre.

Ein ähnliches Beispiel findet sich in Zeile 48 (T3), wo ebenso ein durch die falsche Transkription eines Homonyms entstandener Fehler in der Übersetzung zwar nicht vollständig ausgelöscht, aber verbessert werden konnte. Dabei handelt es sich um „donner“ (T2: Z48) [geben] anstelle von „données“ (T1: Z52) [Daten], welches mit „Daten zu geben“ (T3: Z48) übersetzt wurde. Somit wurde zwar das fälschlicherweise transkribierte Verb mitübersetzt, aber zusätzlich das im Ausgangstext eigentlich gemeinte Wort ‚Daten‘ in die Übersetzung miteingefügt, obwohl dies im Transkript, also der Grundlage für die maschinelle Übersetzung, in dieser Form nicht vorhanden war. Demnach ist wohl wieder davon auszugehen, dass aufgrund des textuellen Kontextes dieser Begriffe die Software die Übersetzung dementsprechend anpassen konnte, um den Inhalt zumindest nur in leicht abgewandelter Form im Vergleich zum mündlichen Ausgangstext wiedergeben zu können.

Ein letztes Beispiel für eine potentielle Verbesserung des Transkriptes im Zuge der maschinellen Übersetzung findet sich im Satz der Zeilen 54-56 (T3). Hier fehlen im Transkript (T2: Z54-56) Satzzeichen, um die Satzgrenzen zu definieren, da es sich ursprünglich um zwei Aussagesätze handelt, die durch einen Fragesatz, der lediglich aus einem Fragewort besteht, voneinander getrennt sind (T1: Z59-61).

Un professeur qui voulait se rendre dans une réserve d'animaux non loin de Shanghai a décidé de refuser d'entrer dans la réserve. Pourquoi ? Parce que les gardiens à l'entrée voulaient scanner son visage. (T1: Z59-61)

un professeur qui voulait se rendre dans un réserve d'animaux non loin de Shanghai à décider de refuser d'entrer dans la réserve pourquoi parce que les gardiens à l'entrée voulaient scanner son visage. (T2 : Z54-56)

ein Professor, der in ein Tierreservat unweit von Shanghai wollte, verweigerte den Zutritt zum Reservat, weil die Wachen am Eingang sein Gesicht scannen sollten. (T3: Z54-56)

Dementsprechend erscheinen diese drei Sätze des mündlichen Ausgangstextes im Transkript als ein zusammenhängender Satz, während in der Übersetzung das Fragewort ‚pourquoi‘ [warum] nicht mehr aufscheint und die beiden übrigen Sätze logisch kohärent miteinander verbunden werden, wodurch ein vollständiger, verständlicher und grammatikalisch korrekter Satz entsteht, der so im zugrundeliegenden Transkript nicht vorhanden ist. An dieser Stelle fand also im Zuge der maschinellen Übersetzung abermals eine Verbesserung des erstellten schriftlichen Ausgangstextes statt, sodass im Zieltext der Sinn und Zweck des ursprünglich mündlich artikulierten Ausgangstextes wiedergegeben werden konnte.

5.2.7 Zusammenfassung

Zusammenfassend lässt sich zum Schritt der maschinellen Übersetzung festhalten, dass die größten Probleme beziehungsweise Unstimmigkeiten, die sich beim Betrachten des deutschsprachigen Zieltextes erkennen lassen, auf Fehlern basieren, die bereits während des Schrittes der Spracherkennung passiert sind, und es sich somit direkt darauf zurückführen lassen. Dabei handelt es sich vor allem um falsch transkribierte Homonyme und falsch oder nicht gesetzte Satzzeichen, die fehlerhafte logische Verknüpfungen, ungrammatische Sätze oder ganze Sinnverfälschungen im Zieltext verursachen. Dementsprechend kann wohl davon ausgegangen werden, dass diese Problematiken nicht an der eigentlichen Funktionsweise der maschinellen Übersetzung liegen. Direkt an der maschinellen Übersetzung selbst lassen sich jedoch auch einige kleinere Problematiken feststellen, die sich vor allem durch eher unpassende oder ungewöhnliche Formulierungen oder zu wörtliche Übersetzungen, die im Deutschen andere Konnotationen als im Französischen haben, äußern aber keine groben Sinnveränderungen bewirken, sondern den Text vermutlich für die Rezipient:innen lediglich als maschinelle Übersetzung erkennbar machen.

Hierbei sei zusätzlich – wie auch im vorangegangenen Kapitel bezüglich der Spracherkennung erläutert wurde – angemerkt, dass die Tatsache, dass es sich um eine recht idealisierte Version eines Ausgangstextes handelt, auch eine Rolle für die Qualität der Übersetzung spielt. So ist bei diesem Text etwa recht wenig zusätzliches situatives oder kulturelles Wissen erforderlich, um den Zieltext adäquat gestalten zu können. Demnach sind auch keine drastischen Änderungen bezüglich des Stils oder Ähnlichem notwendig, wie sich auch später bei der Betrachtung der Zieltextproduktionen der Humandolmetscher:innen zeigen wird. Auch finden sich im Ausgangstext keine lexikalischen oder syntaktischen Ambiguitäten, die für eine gelungene maschinelle Übersetzung eine Schwierigkeit darstellen könnten.

5.3 Sprachausgabe

Im Rahmen dieses Kapitels erfolgt nun die Analyse des dritten und letzten Schrittes der maschinellen Dolmetschung – der Sprachausgabe. Diese soll – wie auch schon die beiden ersten Schritte – zunächst einzeln und unabhängig von den anderen beiden analysiert werden, um in weiterer Folge besser erkennen zu können, an welchen Stellen im Prozess der maschinellen Dolmetschung welche Fehler beziehungsweise Störfaktoren oder Problematiken zu beobachten sind.

5.3.1 Ambiguitäten, Abkürzungen und Fremdwörter

Zunächst lässt sich feststellen, dass in dem deutschen Zieltext der maschinellen Übersetzung und damit jenem Text, der der Sprachausgabe zu Grunde liegt, keine Homographie oder syntaktische Ambiguitäten vorkommen, wodurch eine potentielle Schwierigkeit oder Fehlerquelle, die in Kapitel 3.3 genauer erläutert wurde, bereits ausgeschlossen werden kann. Weiters kommen auch keine Fremdwörter oder Akronyme vor, für die Ähnliches gelten würde.

Einzig eine Abkürzung lässt sich im Text finden, wobei hier anzumerken ist, dass diese Abkürzung durch die maschinelle Übersetzung entstanden ist und so im Ausgangstext ursprünglich gar nicht vorhanden war. Dabei handelt es sich um „d. h.“ (T3: Z48), also die Abkürzung für ‚das heißt‘, welche im Normalfall nur in der Schriftsprache vorkommt und im mündlichen Gebrauch in der Langform ausgesprochen wird. Dies gilt dabei auch für das

Französische, weshalb die Phrase „c’est-à-dire“ (T2: Z48) auch hier nicht in der abgekürzten Version artikuliert wurde. Im deutschen Zieltext, der durch die Software der maschinellen Übersetzung produziert wurde, findet sich hingegen die abgekürzte Variante, was in der Schriftsprache auch adäquat wäre. Nun wird dies durch die Sprachausgabe allerdings nicht als vollständige Phrase, sondern als einzelne Buchstaben inklusive der gesetzten Punkte ausgesprochen, sodass es für Rezipient:innen des gesprochenen Zieltextes, die die Schriftform nicht vorliegen haben, wohl eher unverständlich wäre, dies so zu hören. Ansonsten lassen sich allerdings keinerlei Aussprachefehler feststellen, die zu Missverständnissen oder Verständnisproblemen führen könnten.

5.3.2 Stimme

Hinsichtlich des Klanges der Stimme ist anzumerken, dass diese recht künstlich klingt und somit eindeutig als synthetische Stimme zu erkennen ist und sich deutlich von einer natürlichen menschlichen Stimme unterscheidet. Dies kann dazu führen, dass es für Rezipient:innen anstrengender ist, dem Inhalt zu folgen, worauf im folgenden Kapitel noch genauer eingegangen werden soll.

Auch bezüglich der Prosodie lassen sich einige Mängel feststellen. Zwischen einzelnen Sätzen oder Absätzen etwa werden oft nur recht kurze Pausen gemacht, während an anderen unpassenden Stellen, wie etwa inmitten eines Satzes oder bei Beistrichen, übermäßig lange Pausen mit einer Länge von ein bis zwei Sekunden vorkommen. Eine solche überlange Pause im Satz lässt sich beispielsweise an folgender Stelle erkennen: „[es ist] in China tatsächlich selten, • dass personenbezogene Daten von Personen mit bösen Absichten weitergegeben werden, [...].“ (T4: Z43-45) An anderer Stelle erfolgt hingegen keine Pause nach der Satzgrenze: „Warum also sind die chinesischen Bürger besonders besorgt darüber? Maßnahme und vor allem weil sie befürchten, • dass [...].“ (T4: Z41-42)

Hierbei gilt es aber auch zu beachten, dass das Problem der Zeichensetzung, die eine wichtige Grundlage für die Prosodie und vor allem das Pausensetzen beim Sprechen darstellt, bereits bei der Erstellung des französischsprachigen Transkriptes durch Google Übersetzer entstanden ist und in weiterer Folge auch auf die deutschsprachige Übersetzung übertragen wurde, sodass viele Satzgrenzen durch fehlerhaft gesetzte oder völlig fehlende Satzzeichen für die Software zur Sprachausgabe wohl nicht eindeutig zu erkennen sind.

Außerdem lässt sich auch bei der Betonung einzelner Wörter festhalten, dass diese recht unnatürlich wirkt, da einzelne Silben vor allem bei langen mehrsilbigen Wörtern abgehackt artikuliert werden oder die Wortbetonung falsch gesetzt wird. Dies betrifft beispielweise Wörter wie „Gesichterkennungstechnologie“ (T4: Z3), „kompliziertes“ (T4: Z9) oder auch „Mobiltelefon“ (T4: Z32). Hinzu kommt, dass die synthetische Stimme sehr monoton spricht und dabei die wenigen vorhandenen Betonungen ebenso unnatürlich klingen und somit unpassend wirken. Dies gilt dabei sowohl für die Wortbetonung als auch die Satzmelodie und die hierzu gehörende Satzbetonung, wodurch abermals das Zuhören und Verstehen für die Rezipient:innen des mündlichen Zieltextes deutlich erschwert werden kann.

5.3.3 Zusammenfassung

Zusammenfassend lässt sich bei der Analyse der Funktionsweise der Sprachausgabe von Google Übersetzer feststellen, dass nahezu keine Problematiken hinsichtlich der Verständlichkeit des Inhaltes aufkommen, wobei anzumerken ist, dass in dem der Sprachausgabe vorliegenden Text keine Homographie beziehungsweise syntaktische Ambiguitäten vorkommen, deren Bedeutung abhängig von der Betonung oder Aussprache variieren würde. Zusätzlich kommen keinerlei Fremdwörter, Akronyme oder Eigennamen vor, wo es ebenfalls zu potentiellen Problemen bei der Aussprache kommen könnte, weswegen der Text eigentlich eine ideale Grundlage darstellt, um von einer Software für Sprachausgabe artikuliert zu werden.

Nichtsdestotrotz ist festzuhalten, dass der Klang der Stimme dieser Software sehr künstlich wirkt und einer menschlichen Stimme nahezu gar nicht gleicht, was Rezipient:innen als potentiellen Störfaktor wahrnehmen könnten. Hinzu kommt das monotone Sprechen ohne starke Satz- oder Wortbetonung, was zusätzlich zu einer schnelleren Ermüdung der Zuhörer:innen beitragen und das Verständnis des Inhalts erschweren kann, worauf im folgenden Kapitel im Rahmen der Bewertung der gesamten Dolmetschung noch genauer eingegangen werden soll.

5.4 Bewertung der maschinellen Dolmetschung

In diesem Kapitel erfolgt nun die Bewertung der gesamten Dolmetschleistung von Google Übersetzer. Das bedeutet, dass die drei einzelnen Schritte, deren jeweilige Auffälligkeiten in den vorangegangenen Kapiteln erläutert wurden, nun als Gesamtprodukt betrachtet werden, wie dies auch bei einer Dolmetschung durch Humandolmetscher:innen der Fall wäre. Dabei soll anhand des in Kapitel 2.3 formulierten Bewertungsschemas erarbeitet werden, wie die Dolmetschleistung insgesamt beurteilt werden kann, um diese anschließend mit den Bewertungen der Humandolmetschungen, welche im folgenden Kapitel erläutert werden, vergleichen zu können. Demzufolge folgt zunächst die Bewertung der Übertragung des Inhaltes, ehe mit jener der Realisierung der Dolmetschung fortgesetzt wird. Abschließend erfolgt dann eine Zusammenfassung und die daraus resultierende insgesamte Bewertung der maschinellen Dolmetschung von Google Übersetzer.

5.4.1 Übertragung des Inhalts

5.4.1.1 Treue zum Ausgangstext

Einer der wichtigsten Punkte in Bezug auf die Bewertung einer Dolmetschung stellt die Treue des Zieltextes zum Ausgangstext dar, welche angibt, in welchem Maße die Wirkung und der Zweck adäquat übermittelt werden können und inwiefern ‚information gains‘ oder ‚information loss‘ vorliegen. Die zwei wohl größten Sinnveränderungen im Vergleich zum Ausgangstext finden sich in Zeile 25 (T3), wo anstelle der eigentlichen Übersetzung ‚ungehorsam‘ für ‚dé-sobéissants‘ (T1: Z27) im Zieltext ‚gehorsam‘ (T3: Z25) aufscheint, was sich durch den oben erläuterten Fehler bei der Spracherkennung erklären lässt. Ein ähnliches Problem, was ebenfalls aufgrund der misslungenen Spracherkennungsfunktion entstanden ist, findet sich bei ‚tatsächlich selten‘ (T3: Z42) als Dolmetschung von ‚en effet pas rare‘ (T1: Z46) [tatsächlich nicht selten]. Bei diesen beiden Textstellen kommt es im Rahmen der maschinellen Dolmetschung also zu einer vollständigen Sinnumkehr des Inhalts des Ausgangstextes, sodass im Zieltext ein gegenteiliger Inhalt vermittelt wird. Dies stellt folglich einen groben Verstoß hinsichtlich der Treue zum Ausgangstext und der damit verbundenen angestrebten Übertragung der Wirkung

in den Zieltext dar. Demnach ist dieser Punkt mit besonderer Gewichtung in die Gesamtbewertung der maschinellen Dolmetschung miteinzubeziehen.

Zwei weitere Beispiele, die ‚information loss‘ beziehungsweise eine Sinnverfälschung des Inhaltes des Ausgangstextes darstellen, finden sich in den Zeilen 46 bis 48 (T3) der Übersetzung. Hierbei kam es aufgrund der falschen Transkription von Homophonen zu einer fehlerhaften Übersetzung von „données telles que les adresses ou les numéros de téléphone“ (T1: Z50-51) mit „Adressen oder Telefonnummern anzugeben“ (T3: Z46) anstelle von ‚Daten wie Adressen oder Telefonnummern‘. Ähnliches gilt auch für die Übersetzung „geht es darum, ihnen biometrische Daten zu geben“ (T3: Z47-48), deren entsprechende Stelle im Ausgangstext „ça concerne leurs données euh biometriques“ (T1: Z52) lautet, was etwa mit ‚das betrifft ihre biometrischen Daten‘ übersetzt werden kann. Diese beiden Textpassagen zeigen somit abermals eine Sinnverfälschung aufgrund der fehlerhaften Spracherkennung, wobei hier zu sagen ist, dass es – im Vergleich zu den beiden oben genannten Stellen – zwar zu keiner korrekten Übertragung des Zweckes des Ausgangstextes kommt, aber davon auszugehen ist, dass Rezipient:innen der maschinellen Dolmetschung sich den eigentlich gemeinten Sinn aus dem Kontext erschließen können, da es sich nicht um eine gegenteilige Darstellung des ursprünglichen Inhalts handelt. So erscheinen diese beiden Textpassagen im Zieltext wohl lediglich als etwas inkohärent und könnten schwer verständlich sein, transferieren aber keinen allzu abweichenden Inhalt. Nichtsdestotrotz ist auch dies klarerweise negativ in die Bewertung bezüglich der Treue zum Ausgangstext und damit jener der maschinellen Dolmetschung miteinzubeziehen.

In Zeile 10 (T3) ist ein weiteres Beispiel für eine Sinnverfälschung zu finden. Hier wird im Zieltext davon gesprochen, dass man sich „kein kompliziertes Passwort merken [muss] und es nicht jedes Mal verschlüsseln [muss], wenn [man das] Telefon verwenden möchte“ (T3: Z10), während im Originaltext die Rede davon ist, das Passwort nicht mehr eingeben zu müssen, um das Telefon entsperren zu können. Dementsprechend wurde der Inhalt an dieser Stelle nicht richtig übertragen, was am Schritt der maschinellen Übersetzung zu begründen ist, da das Wort „encoder“ (T2: Z11) [entschlüsseln, *hier*: eingeben] richtig transkribiert, aber für diesen Kontext nicht richtig übersetzt wurde. Allerdings ist davon auszugehen, dass diese Sinnverfälschung nicht zu allzu großen Verständnisproblemen führen sollte, da aus dem Kontext eigentlich klar ist, was diese Textpassage inhaltlich aussagen soll.

Anders sieht es hingegen in der Textpassage der Zeilen 57-64 (T3) im Zieltext aus:

Diese Art von Fall ist in China ziemlich neu und beispiellos und zeigt daher, dass die Bürger in China die Nase voll haben, aber was die chinesische Regierung dazu treibt, Gesichtserkennung in China durchzusetzen alle Aspekte des täglichen Also offiziell

offensichtlich, wie ich sagte, dies ist eine Maßnahme, um die Sicherheit der Bürger zu gewährleisten, aber Sie sollten auch wissen, dass die chinesische Regierung daran interessiert ist, China zu einem technologischen Führer zu machen, deshalb die Regierung, n zögern Sie nicht, Gesichtserkennung oder künstliche Intelligenz zu unterstützen Programme im Allgemeinen. (T3: Z57-64)

Hierbei handelt es sich um einen nahezu völlig unverständlichen Textabschnitt, da einerseits aufgrund von falsch gesetzten Satzzeichen durch die Spracherkennungssoftware keine beziehungsweise falsche Satzgrenzen vorhanden sind und es andererseits auch zu Problemen bei der maschinellen Übersetzung kam. Dadurch ist dieser Satz nicht nur von grammatikalischen und syntaktischen Fehlern und falschen Verknüpfungen, sondern auch von inhaltlichen Einbußen geprägt, die sich im Vergleich mit dem Ausgangstext erkennen lassen. Somit ist die Textpassage deutlich negativ in die Gesamtbewertung der Dolmetschung miteinzubeziehen, da hier große Informationslücken bestehen und der Text auch auf die Rezipient:innen irritierend wirken könnte. Hierbei ist vor allem zu bedenken, dass, wenn dieser Textabschnitt in geschriebener Form vorliegt, noch grob erkennbar ist, welche Inhalte eigentlich transferiert werden sollen, dies jedoch bei der mündlichen Präsentation nahezu völlig verloren geht. So fehlen zusätzlich jegliche Pausen, Betonungen oder ähnliches, um das Gesagte strukturiert wahrnehmen zu können. Deshalb sind hier ganz klar die Treue zum Ausgangstext und die damit verbundene Übertragung der Wirkung und des Zwecks nicht gegeben, weshalb dies als deutlicher Schwachpunkt in die Gesamtbewertung miteinfließen muss.

Eine letzte Textpassage, die ebenfalls einen nahezu unverständlichen Teil des Zieltextes bildet, findet sich gegen Ende dessen:

[...] es stellt eine wertvolle Datenbank für Unternehmen und Verwaltung dar, um zu versuchen, die Bevölkerung zu beruhigen. Die Regierung hat Richtlinien Nicht verbindlich, um die Erhebung und Verwendung von Daten zu regeln. Aber derzeit hat China keine wirkliche Gesetzgebung zu diesem Anscheinend wird also ein Gesetz entworfen, aber wird es ausreichen, um die Bürger zu chinesische Bürger über die Zukunft erzählen uns. (T3: Z66-71)

Hierbei liegen ebenfalls falsche Satzgrenzen und Verknüpfungen sowie einige Übersetzungs- und Grammatikfehler vor, weswegen auch dieser Abschnitt für Rezipient:innen der maschinellen Dolmetschung insgesamt kaum verständlich ist und sich bei einem lediglich einmaligen Anhören wohl kein wirkliches Erschließen des Inhaltes ergeben kann, wodurch Selbiges wie im eben erläuterten Abschnitt gilt und dies eine negative Beeinflussung der Gesamtbewertung mit sich zieht. Da sich bei dieser Textstelle auch Kohärenzprobleme ergeben, wird darauf auch im nächsten Kapitel eingegangen.

5.4.1.2 Kohärenz

Neben der Treue zum Ausgangstext stellt die Kohärenz der Dolmetschung einen weiteren wichtigen Punkt hinsichtlich deren Bewertung dar. Hierbei lässt sich bei der maschinellen Dolmetschung von Google Übersetzer feststellen, dass die Dolmetschung über weite Teile kohärent ist, jedoch finden sich auch einige Stellen, die logisch nicht schlüssig sind und Widersprüche in sich bergen. Dabei fällt auf, dass sich vor allem gegen Ende der Dolmetschung Fehler dieser Art häufen, die zum größten Teil auf misslungene Spracherkennung und damit einhergehende falsche Zeichensetzung oder die falsche Transkription von Homonymen zurückzuführen sind. Demnach liegt die Vermutung nahe, dass je länger die Spracherkennung funktionieren muss und je länger demnach das Transkript wird, desto mehr die Qualität dessen schwindet.

Konkret gesagt brechen etwa die Zeilen 40-44 (T3) mit der Kohärenz des Textes und stellen für Rezipient:innen eine Verständnishürde dar. Hier tritt nicht nur das bereits thematisierte Problem auf, dass im Zuge der Spracherkennung der Negationspartikel im Ausgangstext nicht erkannt und demnach auch nicht übersetzt werden konnte und demzufolge eine Sinnumkehr zu erkennen ist, sondern auch, dass aufgrund der falschen Satzzeichensetzung die Sätze hier falsch miteinander verknüpft sind. So heißt es im Zieltext:

und vor allem weil sie befürchten, dass ihre so erfassten Daten tatsächlich an andere Unternehmen übermittelt und weiterverkauft werden könnten, ist es in China tatsächlich selten, dass personenbezogene Daten von Personen mit bösen Absichten weitergegeben werden, die personenbezogene Daten von Benutzern an andere Unternehmen weiterverkaufen. (T3: 40-44)

Hier wird also inhaltlich suggeriert, dass der Grund dafür, dass es in China selten ist, dass Daten weiterverkauft werden, darin liegt, dass die Bevölkerung befürchtet, dass diese Daten verkauft werden. Diese Verknüpfung gibt es im Ausgangstext nicht. Es ist vielmehr der Fall, dass diese kausale Verknüpfung mit dem vorangegangenen Satz hergestellt wird, sodass es wie folgt heißt:

Alors, pourquoi est-ce que les citoyens chinois sont particulièrement inquiets de cette mesure ? Bien, surtout parce qu'ils craignent que leurs données ainsi enregistrées puissent être transmises et revendues à d'autres entreprises.¹⁷ (T1: Z43-46)

Dementsprechend handelt es sich bei dieser Textpassage des Zieltextes der maschinellen Dolmetschung um eine Inkohärenz, was zu Verständnisproblemen führen könnte. Ähnliches lässt

¹⁷ Also warum sind die chinesischen Bürger besonders besorgt über diese Maßnahme? Gut, vor allem weil sie befürchten, dass ihre so gespeicherten Daten an andere Unternehmen übermittelt und weiterverkauft werden könnten.

sich auch bei der Verknüpfung der Sätze der Zeilen 66-68 (T3) feststellen. Hier lautet der Zieltext wie folgt:

In China gibt es ungefähr 850 Millionen Mobiltelefonbenutzer, es stellt eine wertvolle Datenbank für Unternehmen und Verwaltung dar, um zu versuchen, die Bevölkerung zu beruhigen. Die Regierung hat Richtlinien Nicht verbindlich [sic], um die Erhebung und Verwendung von Daten zu regeln. (T3: Z66-69)

Im Ausgangstext wird jedoch folgendes gesagt:

Ça constitue une précieuse base de données pour les entreprises et les administrations. Alors, pour tenter de rassurer la population le gouvernement a publié des directives non contraignantes visant à réglementer la collecte et l'utilisation des données.¹⁸ (T1: Z74-77)

Hier wurden folglich im Rahmen der maschinellen Übersetzung abermals falsche Verknüpfungen hergestellt, deren Ursprung sich in der fehlerhaften Spracherkennung und den damit einhergehenden falsch gesetzten Satzzeichen begründen lässt. So wurde also die ursprüngliche inhaltliche Verknüpfung des Ausgangstextes, laut der die Richtlinien veröffentlicht wurden, um die Bevölkerung zu beruhigen, derart umgeändert, dass es nun heißt, dass die Datenbank der Handybenutzer den Zweck hat, die Bevölkerung zu beruhigen. Demnach handelt es sich hier wie beim vorherigen Punkt um eine Abänderung des Inhalts, was wiederum zu Verständnisproblemen bei den Rezipient:innen führen kann, da rein vom Zieltext her nicht klar ist, wie die Beruhigung der Bevölkerung und das Verfassen der Richtlinien miteinander in Verbindung zu bringen sind.

5.4.1.3 Pragmatik und Terminologie

In Bezug auf die Pragmatik ist festzuhalten, dass seitens der Software für maschinelle Übersetzung prinzipiell keine pragmatischen Anpassungen möglich sind, sofern diese im Vorfeld nicht explizit programmiert wurden. Jedenfalls fehlt es einer solchen Software grundsätzlich an situativem Wissen, was dazu führt, dass etwa keine geschlechtergerechte Sprache im deutschen Zieltext verwendet wurde. Außerdem scheint die verwendete Begrüßung „Hallo“ (T3: Z1) für den gegebenen Kontext gerade im Vergleich zum französischen Ausgangstext mit der Begrüßung „bonjour“ (T2: Z1) [guten Tag] zu informell und demnach nicht adäquat angepasst. Hinzu

¹⁸ Dies stellt eine wertvolle Basis an Daten für die Unternehmen und die Verwaltungen dar. Also um zu versuchen, die Bevölkerung zu beruhigen, hat die Regierung unverbindliche Richtlinien veröffentlicht, die das Sammeln und Verwenden von Daten reglementieren sollen.

kommt, dass sich in der Übersetzung das Wort „Führer“ (T3: Z62) als Übersetzung für „leader“ (T2: Z62) wiederfindet, welches im Deutschen jedoch eine deutlich andere Konnotation als das verwendete Wort im Französischen aufweist und demnach in diesem Kontext nicht passend ist und angepasst werden müsste.

Hinsichtlich der verwendeten Terminologie ist anzumerken, dass bis auf „reconnaissance faciale“ (u.a. T2: Z1) [Gesichtserkennung] keine relevanten Termini vorkommen, die für die Bewertung der Qualität der Dolmetschungen von Relevanz sind, weswegen es sich bei diesem Ausgangstext – im Hinblick auf die Terminologie – auch um keinen schwierig zu dolmetschenden Text handelt. Die Dolmetschung dieses Terminus erfolgt im Zieltext konsequent und mit der passenden deutschen Entsprechung „Gesichtserkennung“ (u.a. T3: Z1), sodass in diesem Punkt die Dolmetschung als gelungen zu betrachten ist.

5.4.2 Realisierung der Dolmetschung

Zusätzlich zur eben erarbeiteten Bewertung der Übertragung des Inhaltes im Rahmen der Dolmetschung stellt deren Realisierung die zweite wichtige Grundlage für die Bewertung der Qualität dar. Dazu zählt etwa die Stimme der Dolmetscher:innen. Bei der maschinellen Dolmetschung durch Google Übersetzer kommt klarerweise eine synthetische Stimme zum Einsatz, die in diesem Fall sehr künstlich klingt und demnach eindeutig als künstliche Stimme zu erkennen ist und nicht wie eine menschliche Stimme klingt. Dies ist somit negativ in die Bewertung miteinzubeziehen, da künstlich wirkende Stimmen seitens Nutzer:innen als Störfaktor erachtet werden, wie in Kapitel 3.3 erläutert wurde.

5.4.2.1 Verzögerungslaute, Füllwörter und Sprechtempo

Andererseits werden durch eine synthetische Stimme auch keine Verzögerungslaute produziert, da diese im Rahmen des Prozesses der Spracherkennung zum größten Teil herausgefiltert werden und somit für den Schritt der Sprachausgabe nicht mehr als Ausgangsmaterial zur Verfügung stehen. Hierbei sei allerdings angemerkt, dass dies klarerweise nur funktionieren kann, wenn diese Filterfunktion der Spracherkennungssoftware auch problemlos funktioniert. So können beispielsweise Verzögerungslaute fälschlicherweise auch als andere Lexeme erkannt

und demnach übersetzt werden, was selten auch hier vorgekommen ist, wie in Kapitel 6.1 genauer erläutert wurde.

Ähnliches gilt auch für Füllwörter, die bei dieser maschinellen Dolmetschung ebenso kaum vorkommen. Das liegt in diesem Fall – neben der angesprochenen Filterfunktion – daran, dass schon im Ausgangstext wenige solche Füllwörter artikuliert wurden und diese rein durch eine maschinelle Dolmetschung – im Gegensatz zu einer Dolmetschung durch menschliche Dolmetscher:innen – gar nicht erst produziert werden. Nichtsdestotrotz zeichnet die Abwesenheit von Verzögerungslauten und Füllwörtern eine qualitativ hochwertige Dolmetschleistung aus, weshalb dieser Punkt positiv in die Gesamtbewertung miteinzubeziehen ist.

Ein weiterer Punkt, der im Rahmen der Realisierung einer Dolmetschung zu bewerten ist, ist das Sprechtempo. Dieses ist in dem Fall etwas langsamer als das Original, vergleichbar mit jenem der Humandolmetscher:innen, und gleicht somit einem natürlichen Sprechtempo. Auch der Rhythmus ist gleichmäßig und durch die nicht vorhandenen Pausen, Atemgeräusche, Verzögerungslaute oder Füllwörter wird dieser auch nicht unterbrochen, sondern bleibt über die gesamte Rede hinweg konstant.

In Bezug auf die Intonation und Prosodie der maschinellen Dolmetschung ist zu sagen, dass – wie bereits in Kapitel 3.3 genauer erläutert wurde – vor allem hinsichtlich der Pausen Schwierigkeiten auftreten, da die Pausen unnatürlich wirken und teilweise nicht an den richtigen Stellen realisiert werden, was für die Rezipient:innen etwas irritierend wirken könnte. Außerdem klingt die synthetische Stimme sehr monoton, da abgesehen von den unnatürlichen Pausen auch die Satz- und Wortbetonungen unnatürlich klingen. Dies liegt daran, dass diese teilweise entweder gar nicht oder an den falschen Stellen vorkommen. Ein solches monotones Sprechen gilt, wie bereits erläutert wurde, als Störfaktor einer Dolmetschung, da die Konzentration der Rezipient:innen dadurch negativ beeinflusst werden kann, wodurch auch das inhaltliche Verständnis beeinträchtigt wird. Aus diesem Grund muss der Punkt der Intonation und Prosodie negativ in die Gesamtbewertung miteinberechnet werden.

Ansonsten ist die Aussprache trotz der synthetischen Stimme klar und gut verständlich und es kommt in diesem Text auch zu keinen Aussprache Fehlern, die so schwerwiegend sind, dass diese in weiterer Folge Verständnisprobleme verursachen könnten, was wiederum keine negativen Auswirkungen auf die Gesamtbewertung hat.

5.4.2.3 Grammatik

In Bezug auf die Sprachrichtigkeit lässt sich bei Betrachtung der Grammatik der maschinellen Dolmetschung von Google Übersetzer festhalten, dass es zu einigen Fehlern bezüglich der korrekten Verwendung von anaphorischen Pronomen kommt, was sich mit der Funktionsweise der maschinellen Übersetzung begründen lässt. Dies liegt daran, dass die Verwendung von anaphorischen Pronomen je nach Sprache und den betreffenden Genera der Referenzsubstantive variiert und demnach nicht wörtlich von einer Sprache in die andere übersetzt werden kann, sondern der Kontext mitbedacht werden muss, damit die korrekte Übersetzung solcher Textpassagen möglich ist. So kommt im Zieltext (T3) in den Zeilen 16 und 37 das neutrale Pronomen „es“ vor, wobei dieses eigentlich als Referenz auf die femininen Substantive „Gesichtserkennung“ (T3: Z15) beziehungsweise „Verpflichtung“ (T3: Z34) fungieren sollte und demnach stattdessen das feminine Pronomen ‚sie‘ stehen sollte. Ein weiteres Problem mit einem Pronomen, das einen Grammatikfehler darstellt und sich somit negativ auf die Bewertung der Dolmetschleistung auswirkt, findet sich in den Zeilen 24 und 25 (T3), wo das anaphorische Pronomen ‚sie‘, das sich abermals auf „Gesichtserkennung“ (T3: Z23) bezieht, gänzlich fehlt, wodurch sich ein unvollständiger und ungrammatischer Satz ergibt.

In Zeile 68 (T3) lässt sich ein weiterer Grammatikfehler ausmachen, der sich durch eine fehlerhafte Flexion und syntaktische Position eines Adjektivs begründen lässt. Hierbei lautet die Textpassage „Richtlinien Nicht verbindlich“, wobei die korrekte Flexion und Satzstellung wohl ‚nicht verbindliche Richtlinien‘ beziehungsweise ‚unverbindliche Richtlinien‘ lauten müsste. Dieser Fehler lässt sich damit begründen, dass es bei der Spracherkennung zu einer fehlerhaften Zeichensetzung gekommen ist, weswegen die Nominalphrase „des directives. Non contraignants“ (T2: Z68) [Richtlinien. Nicht verbindlich] durch einen Punkt getrennt ist. Das hat wohl die maschinelle Übersetzung insofern erschwert, dass diese Satztrennung im Zieltext zwar nicht mehr vorhanden ist, aber auch die eigentliche Zusammengehörigkeit der Nominal- und Adjektivphrase nicht richtig übersetzt werden konnte.

5.4.2.4 Syntax

Darüber hinaus lassen sich auch noch einige syntaktische Fehler im Zieltext der maschinellen Dolmetschung ausmachen. So finden sich einige Sätze, in denen die Satzgliedstellung nicht korrekt ist, wodurch sich ungrammatische und teilweise für die Rezipient:innen auch nahezu unverständliche Sätze ergeben. In Zeile 32 (T3) etwa lautet eine Textpassage „um erkannt zu werden zur Gesichtserkennung“, wobei hierbei nicht nur ‚zur Gesichtserkennung‘ fälschlicherweise nach dem Prädikat positioniert ist, sondern auch die Präposition falsch übersetzt wurde, weswegen diese Textpassage insgesamt schwer verständlich ist. Die Ursache dieses Fehlers lässt sich dabei allerdings nicht klar verorten, da diese Textpassage im Zuge der Spracherkennung richtig transkribiert wurde und ausgehend von der Struktur des Satzes im Ausgangstext auch keine zu erwartenden Schwierigkeiten ersichtlich sind, da eine recht wörtliche Übersetzung eigentlich adäquat gewesen wäre.

Ein weiteres Beispiel für eine fehlerhafte Satzstruktur findet sich in den Zeilen 39-41 (T3), wo es wie folgt heißt: „Warum also sind die chinesischen Bürger besonders besorgt darüber? Maßnahme und vor allem weil sie befürchten, dass [...]“ (T3: Z39-41). Hier steht das Wort „Maßnahme“ fälschlicherweise zu Beginn des Satzes, obwohl es eigentlich zum vorherigen Satz gehört und an dessen Ende stehen müsste. Demnach handelt es sich bei diesem Folgesatz um einen ungrammatischen und aufgrund der Nennung von ‚Maßnahme‘ auch nicht wirklich verständlichen Satz. Die Problematik lässt sich hier abermals mit der fehlerhaften Spracherkennung begründen, da im Transkript kein Satzzeichen zwischen diesen beiden Sätzen eingefügt wurde, dies aber im Zuge der maschinellen Übersetzung geschehen ist, sodass es sich bei „Warum also sind die chinesischen Bürger besorgt darüber?“ (T3: Z39-40) um einen fehlerfreien Satz und damit sogar um eine Verbesserung im Vergleich zum ausgangssprachlichen Transkript handelt, in den das Adverb ‚darüber‘ zusätzlich eingefügt wurde, um eben diesen syntaktisch korrekten Satz bilden zu können. Laut Ausgangstext sollte an dieser Stelle allerdings ‚über die Maßnahme‘ stehen, weswegen deren Nennung im Folgesatz redundant und syntaktisch nicht korrekt ist.

Ähnliches gilt auch für die Sätze in den Zeilen 58-64 (T3), in denen ebenfalls die Satzstellung nicht grammatikalisch korrekt ist, was sich abermals damit erklären lässt, dass die Zeichensetzung und die damit verbundenen Satzgrenzen im Zuge der Erstellung des Transkriptes des Ausgangstextes durch die Spracherkennungssoftware nicht fehlerfrei vollzogen wurden:

[...] aber was die chinesische Regierung dazu treibt, Gesichtserkennung in China durchzusetzen alle Aspekte des täglichen Also offiziell, wie ich sagte, dies ist eine

Maßnahme, um die Sicherheit der Bürger zu gewährleisten, aber Sie sollten auch wissen, dass die chinesische Regierung daran interessiert ist, China zu einem technologischen Führer zu machen, deshalb die Regierung, n zögern Sie nicht, Gesichtserkennung oder künstliche Intelligenz zu unterstützen Programme im Allgemeinen. (T3: Z58-64)

Demnach handelt es sich auch hier um Sätze, die nicht nur grammatikalisch falsch sind, was Rezipient:innen irritieren könnte, sondern in denen auch Verknüpfungen zwischen den Sätzen, sowie inhaltliche Zusammenhänge nicht adäquat wiedergegeben werden, wodurch es in weiterer Folge zu Verständnisproblemen und damit nur zu einer eingeschränkten Weitergabe des Inhalts kommt. Selbiges gilt auch für den Satz in den Zeilen 70-71 (T3), wo die Selbstkorrektur „les chito/citoyens chinois“ (T1: Z79) [die chinesischen Bürger] sowie die undeutliche Aussprache von „bien“ (T1: Z79) [gut] zu Problemen bei der Spracherkennung führten. Hinzu kommt, dass auch die Satzzeichen nicht korrekt gesetzt wurden, weswegen es im Zuge der maschinellen Übersetzung zu Problemen kam und im Zieltext ein ungrammatischer und unvollständiger Satz steht, der wohl abermals Verständnisprobleme bei Rezipient:innen hervorrufen würde.

In Zeile 70 (T3) findet sich ein weiteres Beispiel für einen syntaktischen Fehler, wobei wieder nicht ersichtlich ist, wie dieser entstanden ist. Hierbei handelt es sich um einen unvollständigen Satz im Zieltext, der ohne Satzzeichen abschließt, obwohl das Wort danach wie am Satzanfang mit einem Großbuchstaben beginnt: „Aber derzeit hat China keine wirkliche Gesetzgebung zu diesem Anscheinend“ (T3: Z69-70). Im Transkript der Spracherkennung allerdings ist nicht nur der Satz an dieser Stelle vollständig und weitestgehend richtig transkribiert, sodass davon ausgegangen werden kann, dass die maschinelle Übersetzung funktionieren müsste, sondern auch die Zeichensetzung korrekt, sodass die Satzgrenze eigentlich klar erkennbar ist, was den Schritt der maschinellen Übersetzung zusätzlich erleichtern sollte. Dementsprechend handelt es sich hier klarerweise um einen syntaktischen Fehler im Zieltext, der in weiterer Folge wiederum zu inhaltlichen Verständnisproblemen führen kann, dessen Ursache allerdings nicht klar ersichtlich, wohl aber im Zuge der maschinellen Übersetzung zu verorten ist.

5.4.2.5 Lexik und Stil

Zusätzlich beinhaltet der Zieltext der maschinellen Dolmetschung einige lexikalische Unstimmigkeiten, die vor allem auf zu direkten beziehungsweise wörtlichen Übersetzungen des Ausgangstextes beruhen. Dazu zählt beispielsweise die Phrase „[wir sind] oft gezwungen, auf China zu verweisen“ (T3: Z11-12), die als Übersetzung von „on est souvent forcé de faire référence à la Chine“ (T2: Z13) fungiert. Dabei handelt es sich um eine wörtliche Übersetzung des französischen Originals, wobei die Aussage im Deutschen eine andere Konnotation hat, da die Formulierung im Französischen im übertragenen Sinn fungieren kann, während im Deutschen eher der wörtliche Sinn im Vordergrund steht und in diesem Kontext nicht von einem wörtlichen ‚Zwang‘ die Rede sein kann. Dies wird auch bei Betrachtung der Dolmetschleistungen der menschlichen Dolmetscher:innen sichtbar, die allesamt diese Passage so umformuliert haben, sodass im Deutschen eine ähnliche Wirkung wie im Französischen erzielt werden kann.

Ähnliches gilt auch für die Passage „Die Gesichtserkennung ist [...] bequem im täglichen Leben [...] installiert“ (T3: Z15-16), was als wörtliche Übersetzung von „La reconnaissance faciale est en effet déjà confortablement installé dans le quotidien“ (T2: Z16) erscheint. Dabei ergibt sich nun abermals das Problem, dass das direkt übertragene Verb ‚installé‘ im Deutschen andere Konnotationen aufweist und in diesem Kontext nicht verwendet werden kann, ohne dass es seitens der Rezipient:innen als störend empfunden werden könnte. Auch dies wird durch die Humandolmetschungen bestätigt, da es auch hier in allen Fällen zu Umformulierungen gekommen ist, um einen adäquaten deutschen Zieltext zu kreieren. Auch die Übersetzung von „des applications sur téléphone portable“ (T2: Z49-50) mit „Mobiltelefonanwendungen“ wirkt etwas unpassend, da die üblichere Formulierung – wie auch die Humandolmetschungen zeigen – wohl ‚Apps‘ lauten würde.

Ein weiteres Beispiel für eine zu wörtliche Übertragung findet sich in den Zeilen 19-20 (T3) mit „die Gesichtserkennung [...] ist gleichbedeutend mit vielen zusätzliche Vorteilen“ als Übersetzung von „la reconnaissance faciale [...] est synonyme de nombreux bénéfices“ (T2: Z19-20). Auch diese Formulierung wirkt im deutschen Zieltext etwas unpassend und unnatürlich, da sie strukturell zu nahe am Ausgangstext liegt, weswegen diese Passage – wie es bei den Humandolmetscher:innen abermals der Fall war – freier formuliert werden könnte, um im Zieltext adäquat zu sein.

In Bezug auf all diese hier erwähnten lexikalischen Problematiken, die sich bei der maschinellen Dolmetschung ergaben, sei allerdings angemerkt, dass es sich um keine Sinnverfälschungen handelt, sondern lediglich die Formulierungen im Zieltext nicht adäquat erscheinen

und demnach störend oder irritierend wirken könnten, während der eigentliche Inhalt dieser Passagen – wie im vorigen Kapitel erläutert wurde – trotzdem in die Zielsprache transferiert werden konnte.

In Bezug auf den Stil der maschinellen Dolmetschung lässt sich allerdings festhalten, dass dieser – abgesehen von den bereits erläuterten grammatikalischen und lexikalischen Fehlern und Ungenauigkeiten – für den Zweck und Kontext der Dolmetschung angemessen erscheint, da es in diesem Setting generell nicht nötig ist, den Stil des Ausgangstextes zu verändern und dieser somit einfach für den Zieltext übernommen werden kann. Demnach können bei diesem Beispiel keine besonderen stilistischen Fehler hervorgehoben werden, da der Ausgangstext selbst keine für eine Dolmetschung relevanten stilistischen Besonderheiten in sich trägt.

5.4.2.6 Länge

Wie bereits in Kapitel 2.2 angesprochen wurde, liegt die Erwartung an die optimale Länge einer Konsekutivdolmetschung etwa bei 70% der Länge des Originaltextes. Bei der maschinellen Dolmetschung von Google Übersetzer ist diese Zeitspanne allerdings deutlich überschritten worden, sodass die Realisierung des Zieltextes insgesamt 6 Minuten und 24 Sekunden in Anspruch nimmt, was einer Länge von 103% im Vergleich mit dem Ausgangstext entspricht. Damit wird dessen Länge sogar überschritten, was ebenfalls negativ in die Bewertung der Dolmetschleistung miteinzubeziehen ist. Hierbei zeigt sich abermals ein Problem der maschinellen Dolmetschung, da die Software klarerweise nur den gesamten dargebotenen Text transkribieren, übersetzen und mündlich wiedergeben kann und keine Redundanzen oder Ähnliches auslassen kann, da dafür Kontextwissen, pragmatisches und inhaltliches Verständnis fehlen. Dolmetscher:innen hingegen können solches Wissen anwenden, um den Zieltext dementsprechend adäquat zu gestalten und die gewünschten bzw. erwarteten Längenvorgaben einzuhalten. Wie dies konkret aussehen kann, wird im weiteren Verlauf der Masterarbeit genauer erläutert, wenn auf die Bewertung der Dolmetschleistung der Humandolmetscher:innen eingegangen wird.

5.4.3 Quantitative Analyse

In diesem Abschnitt sollen nun die oben ausführlich erläuterten Mängel und Fehler an der Dolmetschung von Google Übersetzer nach denselben Kriterien quantitativ ausgewertet und kategorisiert werden, um einen noch umfangreicheren Vergleich mit den menschlichen Dolmetschleistungen schaffen zu können.

In Bezug auf die Übertragung des Inhalts finden sich insgesamt ein leichter, sieben mittlere und acht schwere Fehler in der gesamten Dolmetschung, die sich auf die Bewertungskategorien Treue zum Ausgangstext, Kohärenz und Pragmatik aufteilen. Hinsichtlich der Treue zum Ausgangstext lassen sich davon ein leichter, vier mittlere und sieben schwere Fehler ausmachen, die sich wie folgt zusammensetzen: Ein leichter Fehler findet sich in Zeile 57 (T3), der sich durch die Verwendung eines unbestimmten Artikels statt des eigentlichen bestimmten Artikels ergibt. Dadurch könnte die Verständlichkeit dieser Textstelle leicht beeinträchtigt sein. Die mittleren Fehler finden sich in den Zeilen 10, 32, 46 und 48 (T3). Bei diesen Fehlern handelt es sich um falsche Verknüpfungen oder die Verwendung von unpassenden Wörtern, was zu einem verminderten Verständnis führen kann. Somit ist die Verständlichkeit der Dolmetschung durch diese Fehler eingeschränkt, aber trotzdem noch in gewissem Maße für die Rezipient:innen möglich, weshalb diese als mittlere Fehler klassifiziert werden. Bei den sieben schweren Fehlern hingegen handelt es sich um vollständige Sinnumkehrungen, wie in den Zeilen 35 und 42 (T3) oder gänzlich unverständliche Passagen der Dolmetschung, wie in den Zeilen 59, 60, 63, 64 und 71 (T3), sodass hier keine Verständlichkeit mehr für die Nutzer:innen gegeben ist.

Bezüglich der Kohärenz der maschinellen Dolmetschung lassen sich zwei schwere Fehler in den Zeilen 67-68 (T3) sowie in den Zeilen 40-44 (T3) ausmachen, da an diesen Stellen falsche Verknüpfung von zwei Satzteilen zu finden ist, wodurch der Sinn dieses Textabschnittes für die Rezipient:innen nicht mehr klar verständlich ist, und die Übertragung des Inhalts somit nicht funktioniert. In Bezug auf die Pragmatik finden sich zwei mittlere Fehler in den Zeilen 1 (T3) aufgrund der fehlenden Begrüßung und 62 (T3) wegen der Verwendung des an dieser Stelle unpassenden Wortes ‚Führer‘ im Deutschen.

Hinsichtlich der Grammatik finden sich ein leichter Fehler durch ein fehlendes Subjekt in den Zeilen 4-25 (T3), der das Verständnis kaum negativ beeinflusst, sowie drei mittlere Fehler (vgl. T30: Z16, Z37, Z70) durch falsche Pronomen, die als Referenz genutzt wurden beziehungsweise ein fehlendes Objekt, durch die das Verständnis etwas eingeschränkt wird. Zusätzlich lässt sich auch ein schwerer Fehler ausmachen, der durch die falsche Flexion und Stellung

eines Adjektivs in Zeile 68 (T3) entsteht, wodurch die Verständlichkeit stark eingeschränkt wird.

Bezüglich der Lexik lassen sich drei leichte Fehler ausmachen, die auf leicht anderen Konnotationen (vgl. T3: Z20, Z50) beziehungsweise einer ungrammatisch wirkenden Wortwahl (vgl. T3: Z27-28) beruhen. Dazu kommen zwei mittlere Fehler (vgl. T3: Z11-12, 15-16), deren Grundlage ebenfalls andere Konnotationen der gewählten Wörter sind. An diesen Stellen wird durch die verwendeten Worte die Verständlichkeit und damit die Qualität der Dolmetschung allerdings stärker negativ beeinflusst als bei den zuvor genannten leichten Fehlern, weswegen diese Fehler hier als mittlere Fehler klassifiziert wurden.

Hauptkategorie	Fehlerkategorie	leicht	mittel	schwer	Gesamt
Übertragung des Inhalts	Treue zum Ausgangstext	1	4	7	12
	Kohärenz	0	0	2	2
	Pragmatik	0	2	0	2
Realisierung der Dolmetschung	Grammatik	1	3	1	5
	Lexik	3	2	0	5
Gesamt		5	11	10	26

Fehleranalyse Google Übersetzer

Insgesamt finden sich folglich in der Dolmetschung von Google Übersetzer – wie auch in der obenstehenden Tabelle erkennbar ist – fünf leichte, elf mittlere und zehn schwere Fehler, wobei klarerweise vor allem die mittleren und schweren Fehler hervorstechen, da diese einen großen negativen Einfluss auf die Verständlichkeit, die Erfüllung der Erwartungen der Nutzer:innen und auf die insgesamt Qualität der Dolmetschleistungen mit sich bringen. Im Rahmen der Bewertung der Humandolmetscher:innen im folgenden Kapitel werden diese ebenfalls nicht nur qualitativ, sondern auch quantitativ analysiert, sodass anschließend in Kapitel 6 auch ein Vergleich dieser quantitativen Analyse der Dolmetschung von Google Übersetzer mit jenen der Humandolmetscher:innen erfolgen kann, um die hier vorliegenden Ergebnisse in Relation setzen zu können.

5.4.4 Gesamtbewertung

Zusammenfassend lässt sich hinsichtlich der Bewertung der maschinellen Dolmetschleistung von Google Übersetzer festhalten, dass sich einige negative Punkte ausfindig machen lassen. So finden sich in der Dolmetschung etwa einige Sinnverfälschungen im Vergleich zum Ausgangstext, sodass die Treue und die damit verbundene Übertragung der Wirkung und des Zwecks des Ausgangstextes nicht über den gesamten Text hinweg gegeben ist. Hierbei fällt vor allem auf, dass gegen Ende des Textes teilweise völlig unverständliche Passagen auftauchen, die vor allem Fehlern in der Spracherkennung geschuldet sind, wozu einige Übersetzungsfehler kommen. Hinzu kommt, dass ebenso aufgrund fehlerhafter Spracherkennung einige Aussagen völlig gegensätzlich gedolmetscht wurden. Demnach kann insgesamt zur Treue festgehalten werden, dass diese über weite Teile des Textes gegeben ist, abschnittsweise allerdings so gravierende Fehler auftreten, die das inhaltliche Verständnis derart stören, dass hier starke Abstriche hinsichtlich der Qualität gemacht werden müssen. Die Kohärenz ist hingegen weitestgehend gegeben, sodass Verknüpfungen innerhalb des Textes meist logisch erscheinen und mit jenen des Ausgangstextes übereinstimmen, jedoch kommt es auch hier aufgrund der Spracherkennung zu einigen kleineren Problematiken und falsch miteinander in Verbindung gebrachten Aussagen, die die Übertragung des Inhaltes ebenso leicht stören.

Da der Software für maschinelles Dolmetschen situatives Wissen fehlt, sind grundsätzlich keine pragmatischen Anpassungen möglich, wobei dies bei dem hier untersuchten Text keine allzu große Rolle spielt, da es abgesehen vom fehlenden Gendern im Zieltext, der eventuell zu informellen Begrüßung zu Beginn der Rede sowie der Verwendung des Wortes ‚Führer‘ zu keinen größeren Problemen durch pragmatische Abweichungen kommt. Ähnliches gilt auch für die Terminologie, da der Text keine auffallenden terminologischen Schwierigkeiten in sich birgt. Der wichtige Terminus der ‚Gesichtserkennung‘ wird dabei konsequent im gesamten Zieltext verwendet.

Hinsichtlich der Realisierung der Dolmetschung ist anzumerken, dass der Klang der Stimme einerseits stark künstlich wirkt, wobei andererseits keine Füllwörter oder Verzögerungslaute vorkommen, weshalb der Rhythmus als gleichmäßig zu erachten ist. Außerdem handelt es sich bei der Zieltextproduktion um ein angenehmes Sprechtempo, das aber etwas langsamer als das Original ist, was aufgrund der fehlenden Kürzungen dazu führt, dass die Zieltextproduktion länger dauert als jene des Ausgangstextes, was nicht den Ansprüchen an eine hochwertige Konsekutivdolmetschleistung entspricht. Auch die Aussprache ist klar und verständlich. In Bezug auf die Prosodie lässt sich allerdings feststellen, dass recht monoton gesprochen

wird und die Betonung einzelner Silben oder Wörter unnatürlich wirkt, wozu kommt, dass auch das Setzen von Pausen nicht den eigentlichen Satzzeichen entspricht und dementsprechend ebenso störend wirken kann.

In Bezug auf Grammatik, Syntax, Lexik und Stil ist schlussendlich festzuhalten, dass es hierbei sowohl aufgrund der Spracherkennung als auch der maschinellen Übersetzung zu diversen kleineren und größeren Fehlern kommt, sodass etwa falsche Pronomen als Referenz angeführt werden und nicht adäquate Lexeme oder unvollständige oder syntaktisch inkorrekte Sätze vorkommen. Dies kann folglich zu einigen Verständnisproblemen beziehungsweise Irritationen bei den Rezipient:innen führen.

Hierbei sei auch hervorgehoben, dass je länger die Dolmetschung dauert, desto mehr Fehler auftauchen. So finden sich zu Beginn des Zieltextes weitaus weniger Fehler als dies gegen Ende des Zieltextes der Fall ist. Auch die nahezu vollends unverständlichen Passagen finden sich allesamt gegen Ende der Dolmetschung. Demnach stellt sich die Frage, ob es sich dabei um einen reinen Zufall handelt, da der Ausgangstext keine Schwierigkeitsunterschiede im Verlauf aufweist, wie auch im weiteren Verlauf der Arbeit anhand den menschlichen Dolmetschleistungen gezeigt werden kann. Andererseits könnte das Programm Google Übersetzer nur für kürzere Abschnitte ausgelegt sein, da möglicherweise nicht genug rechnerische Kapazitäten vorhanden sind. Dabei handelt es sich jedoch vorerst rein um Spekulationen, da diesem Phänomen im Rahmen dieser Arbeit nicht genauer nachgegangen werden kann.

Insgesamt handelt es sich bei der maschinellen Dolmetschung von Google Übersetzer nun eine teilweise adäquate Dolmetschung, wobei vor allem hinsichtlich der Realisierung der Dolmetschleistung deutliche Abstriche gemacht werden müssen und auch die Treue des Zieltextes zum Ausgangstext aufgrund von teils gegensätzlichen Aussagen nicht immer gegeben ist, wobei dies für eine endgültige Bewertung noch in Relation mit der Leistung der menschlichen Dolmetscher:innen gesetzt werden muss. In den folgenden Kapiteln soll – nach der quantitativen Analyse der Dolmetschung von Google Übersetzer – demzufolge auf die Bewertung der Leistungen der Humandolmetscher:innen eingegangen werden, mit denen jene hier behandelte Bewertung der Dolmetschleistung von Google Übersetzer verglichen werden soll, um abschließend die Möglichkeiten und Grenzen des maschinellen Dolmetschens mit dieser App eruieren zu können.

5.5 Bewertung der Humandolmetscher:innen

Im Rahmen dieses Kapitels soll nun die Bewertung der Leistungen der Humandolmetscher:innen erfolgen, damit eine Vergleichsbasis geschaffen werden kann, mit welcher die Bewertung der maschinellen Dolmetschung von Google Übersetzer in Relation gesetzt werden kann. Für die Analyse werden – wie im Zuge der Methodenbeschreibung in Kapitel 4 erläutert wurde – vier Dolmetschleistungen herangezogen, die hier – wie auch zuvor die maschinelle Dolmetschung – nach dem in Kapitel 2 erarbeiteten Schema analysiert und bewertet werden sollen.

5.5.1 Übertragung des Inhalts

Hinsichtlich der Treue zum Ausgangstext und der korrekten Übertragung von Wirkung und Zweck in den Zieltext lässt sich bei Betrachtung der vier Dolmetschleistungen der Humandolmetscher:innen festhalten, dass die insgesamt Wirkung des Originaltextes bei allen vier in die Zielsprache übertragen werden konnte, da die wichtigsten Informationen des Ausgangstextes in den Zieltexten ebenso aufzufinden sind, keine Widersprüche vorkommen und somit der Zweck der Dolmetschungen als erfüllt zu erachten ist. Allerdings finden sich in den Dolmetschungen auch einige kleinere Ungenauigkeiten bezüglich der Treue, die bei einer genaueren Analyse sichtbar werden und im Folgenden genauer erläutert werden sollen.

5.5.1.1 Auslassungen

Neben diesen Kürzungen von Redundanzen des Ausgangstextes, die keine negativen Auswirkungen auf die Übertragung des Inhaltes in den Zieltext haben, finden sich in den Dolmetschungen allerdings auch einige Kürzungen im Vergleich zum Original, die durchaus mit einigen inhaltlichen Abstrichen oder der Auslassung von mehr oder weniger relevanten Details einhergehen und somit als ‚information loss‘ zu bewerten sind.

Ein Beispiel für solche Auslassungen findet sich bei der Dolmetschung von Aufzählungen, wenn einzelne Punkte der Ausgangsrede weggelassen werden. So zum Beispiel bei der Textstelle „ça concerne leurs données biométriques. C’est-à-dire, leurs empreintes digitales, les

scans de leurs visages ou encore de l'iris de leur œil“¹⁹ (T1: Z52-53). In drei der Dolmetschungen wurden an dieser Stelle lediglich der „digitale[n] Fingerabdruck (T6: Z45-46; 7: Z48) bzw. „Fingerabdrücke“ (T5: Z45), sowie die „Iris“ (T5: Z45; T6: Z46; T7: Z48) gedolmetscht und „die Gesichtszüge“ (T8: Z46), wie es in der vierten Dolmetschung vorkommt, weggelassen. Bei einem solchen Auslassen von einzelnen Punkten einer Aufzählung, sodass der Inhalt trotzdem ohne größere Einbußen transferiert werden kann, handelt es sich an dieser Stelle um eine legitime Dolmetschstrategie, da für die Nutzer:innen insgesamt kein großer ‚information loss‘ besteht und diese Stelle in den Zieltexten flüssig und kohärent erscheint.

Eine weitere Auslassung betrifft die folgende Textstelle :

Mais il faut également savoir que le gouvernement chinois est désireux de faire de la Chine un leader technologique. Par conséquent, le gouvernement n'hésite pas à offrir son soutien à des programmes de reconnaissance faciale ou d'intelligence artificielle en général.²⁰ (T1 : Z67-71)

Hier wurde in drei Dolmetschungen ausgelassen, dass China „Firmen, die sich mit Gesichtserkennungsprogrammen [...] beschäftigen, Unterstützung an[bietet]“ (T8: Z62) und nur erwähnt, dass China eine führende Rolle in Bezug auf Gesichtserkennung und künstliche Intelligenz innehaben möchte (vgl. T5: Z54-56; T6: Z59) bzw. dass China „jede Gelegenheit aus[nutzt]“ (T7: Z61), um diese Führungsrolle zu übernehmen. Hier handelt es sich folglich um einen gewissen ‚information loss‘, der jedoch das insgesamt Verständnis und die Kohärenz dieses Abschnittes nicht negativ beeinflussen sollte.

Weitere ähnliche Fälle solcher kleineren inhaltlichen Auslassungen, die jedoch das Gesamtverständnis der betreffenden Textpassagen und die Kohärenz nicht negativ beeinflussen, sondern nur inhaltliche Details weglassen, finden sich in folgenden Abschnitten der Dolmetschungen: In dem Textabschnitt „la Chine ne dispose pas de vraies législations à ce sujet. Alors, apparemment, une loi est en cours de rédaction.“²¹ (T1: Z77-78) wird in einer Dolmetschung gekürzt, dass an einem „Gesetz [...] gearbeitet [wird] (T7: Z66), sodass lediglich folgendes gedolmetscht wurde: „aber bis jetzt gibt es nicht wirklich die Gesetzgebung, die nötig wäre dazu.“ (T5: Z61-32). Ähnliches gilt für das Kürzen der Textstellen „Elle permet notamment d'accélérer les contrôles d'identité dans les transports. Elle permet aussi d'identifier des

¹⁹ Das betrifft ihre biometrischen Daten. Das heißt, ihre digitalen Fingerabdrücke, Scans ihrer Gesichter oder auch der Iris ihrer Augen.

²⁰ Aber man sollte auch wissen, dass die chinesische Regierung aus China gerne einen technologischen Leader machen möchte. Folglich zögert die Regierung nicht, ihre Unterstützung zu Programmen zur Gesichtserkennung oder künstlicher Intelligenz im Allgemeinen anzubieten.

²¹ China verfügt über keine wirklichen Gesetzgebungen zu diesem Thema. Nun ist offensichtlich ein Gesetz dabei verfasst zu werden.

citoyens désobéissants.“²² (T1: Z26-27) und „Elle a suscité l’inquiétude voire la révolte de certains citoyens“²³ (T1: Z29) in Dolmetschung 7. Hier heißt es lediglich, dass die Gesichtserkennung benutzt wird, „um eventuelle problematische Bürger identifizieren zu können“ (T7: Z27) und, dass „diese Maßnahme [...] schon für viele Sorge [bereitet].“ (T7: Z29). Demzufolge wurden die Zeitersparnis und das Hervorrufen von Protesten ausgelassen, sodass auch hier inhaltliche Aspekte verloren gehen, die jedoch für die Rezipient:innen des Zieltextes nicht sichtbar sind und auch keine allzu wichtigen Schlüsselinformationen beinhalten.

5.5.1.2 Sinnverfälschungen

Beispielsweise kann eine solche Sinnverfälschung auf dem Weglassen von Informationen beruhen, wie in „diese Daten sammeln“ (T5: Z46) bzw. „Daten [...] speichern“ (T6: Z47-48), während diese Stelle im Ausgangstext „[elles] étaient soupçonnées de collecter [...] les données“²⁴ (T1: Z54-56) lautet. Somit wird der ursprüngliche Inhalt des Verdachtes nicht in die Zieltexte transferiert, sondern dort als Fakt dargestellt. Dementsprechend wird zwar der Inhalt des Sammelns von Daten für die Rezipient:innen verständlich, nicht jedoch, dass dies eigentlich nur als Verdacht geäußert wurde, sodass es sich hier um eine Sinnverfälschung handelt.

Eine weitere Art, wie eine Sinnverfälschung auftreten kann, sind Verallgemeinerungen oder Spezifizierungen. Hierbei kann es sich in gewissem Maße auch um eine Dolmetschstrategie handeln, jedoch sollte stets die Übertragung des Inhalts gewährleistet sein. In einer Dolmetschung ist beispielsweise von „U-Bahnen“ (T6: Z24) die Rede, obwohl es im Ausgangstext „les transports“ (T1: Z27) [Transportmittel] heißt. Somit wurde aus dem allgemeineren Begriff der Transportmittel der Spezifischere der U-Bahn gewählt. Aus dem Kontext heraus kann davon ausgegangen werden, dass möglicherweise auch U-Bahnen dazuzählen, jedoch sicherlich auch weitere Verkehrsmittel, die an dieser Stelle in der Dolmetschung nicht erwähnt werden.

Ein solches Beispiel für eine Verallgemeinerung findet sich bei der folgenden Textstelle:

Der offizielle Grund dafür ist die Sicherheit in öffentlichen Räumen, aber eigentlich möchte China auch an der Spitze der technologischen Entwicklung in der Welt sein und

²² Sie erlaubt besonders, die Identitätskontrollen in den öffentlichen Verkehrsmitteln zu beschleunigen. Sie erlaubt es auch, ungehorsame Bürger zu identifizieren.

²³ Sie hat auch Sorge und sogar Aufstände von bestimmten Bürgern hervorgerufen.

²⁴ [sie] wurden verdächtigt, Daten zu sammeln

das heißt, sie [Anm. chinesische Regierung] nutzen quasi jede Gelegenheit aus, um dies zu tun. (T7: Z59-61)

Im Ausgangstext heißt es jedoch:

Mais il faut savoir également que le gouvernement chinois est désireux de faire de la Chine un leader technologique. Par conséquent, le gouvernement n'hésite pas à offrir son soutien à des programmes de reconnaissance faciale ou d'intelligence artificielle en général.²⁵ (T1: Z67-71)

Demnach wurde im Ausgangstext genauer erläutert, was die chinesische Regierung plant, um die technologische Entwicklung rund um künstliche Intelligenz und Gesichtserkennung voranzutreiben, während dies in der Dolmetschung stark generalisiert wird. So wird zwar der Inhalt transferiert, dass China sich dahingehend weiterentwickeln möchte, aber nicht wie genau dies passieren soll. Bei dieser Generalisierung bzw. der Spezifizierung wird demnach zwar der grundlegende Inhalt in die Dolmetschung mit aufgenommen, jedoch können wichtige Bestandteile des Sinnes nicht transferiert werden, sodass es sich hier um leichte Sinnverfälschungen handelt.

Ein weiterer Punkt, wie der Sinn des Ausgangstextes verfälscht werden kann, ist die fehlerhafte Dolmetschung von einzelnen Lexemen. Hierzu zählt zum Beispiel die Verwendung von „Kästchen in den Supermärkten“ (T7: Z18) anstelle von „Supermarktkassen“ (T5: 17; T6: Z16) als Übersetzung für „caisses aux supermarchés“ (T1: Z18-19). Folglich könnte es hier zu Verständnisproblemen kommen, da keine adäquate Übersetzung gewählt wurde. Ähnliches gilt für die Verwendung von „Naturreserve“ (T7: Z53, 56) als Übersetzung von „une réserve d'animaux“ (T1: Z59-60, 63), wobei die passendere Übersetzung etwa „Reservat“ (T5: Z49) lauten würde. Auch hier könnte es aufgrund der unklaren Bedeutung des gewählten Wortes zu Verständnisproblemen kommen, wodurch der Sinn des Ausgangstext nicht in den Zieltext transferiert werden kann. Selbiges gilt für „Einkäufe machen“ (T7: Z19) als Dolmetschung von „payer ses courses“ (T1: Z19). Bei Betrachtung all dieser Beispiele ist festzuhalten, dass mit Einbeziehung des Kontextes für die Rezipient:innen klar sein könnte, welchen Sinn der Zieltext eigentlich haben soll, jedoch kann es trotzdem zu Verständnisproblemen kommen. Demnach sind hier Sinnveränderungen nicht auszuschließen, sondern es ist eher anzunehmen, dass es zu Sinnveränderungen kommt, was einen negativen Einfluss auf die Gesamtbewertung der Dolmetschung hat.

²⁵ Aber man muss auch wissen, dass die chinesische Regierung gewillt ist, aus China einen technologischen Führer zu machen. Folglich zögert die Regierung nicht, ihre Unterstützung für Programme zur Gesichtserkennung oder künstlicher Intelligenz im Allgemeinen anzubieten.

Hinzu kommen weitere Beispiele aus den Dolmetschungen, in denen Sinneinheiten verändert oder sogar konträr in Bezug auf den Ausgangstext wiedergegeben werden. In einer Dolmetschung wird von einem „digitalen Fußabdruck“ (T8: Z46) gesprochen, wobei eigentlich es eigentlich um „empreintes digitales“ (T1: Z53) also um „Fingerabdrücke“ (T5: Z45) geht, wodurch der ursprüngliche Sinn des Ausgangstextes durch die Dolmetschung verändert wurde. Ähnliches gilt für einen Textausschnitt einer Dolmetschung, in dem erläutert wird, dass es „etwa 850 Millionen Menschen [gibt], die Handys nutzen“ (T8: Z64), wobei im Original Folgendes gesagt wird: „En Chine, on compte environ 850 millions d'utilisateurs de téléphones portables.“²⁶ (T1: Z73-74). Demnach wurde in der Dolmetschung ausgelassen, dass sich diese Zahl nur auf die Menschen in China bezieht, sodass es weltweit klarerweise deutlich mehr Menschen sein müssten. Somit handelt es sich hier um eine Sinnverfälschung, gleichwohl es in dem gesamten Text um Menschen in China geht, sodass aus dem Kontext der eigentliche Sinn erschlossen werden könnte.

In zwei Dolmetschungen findet sich allerdings auch eine Textstelle, die inhaltlich das Gegenteil im Vergleich zum französischen Original aussagt. Diese Stelle lautet im Ausgangstext wie folgt: „Ce genre de cas est tout à fait nouveau et sans précédent en Chine.“²⁷ (T1: Z63-64). In einer der Dolmetschungen heißt es jedoch „das ist nur ein Beispiel, das ist nichts Neues.“ (T5: Z51) und in einer weiteren wie folgt: „und das ist eines der Beispiele“ (T6: Z55-56). In diesen beiden Fällen kam es also zu einer völligen Sinnumkehr, sodass für die Rezipient:innen nicht klar wird, dass es sich bei dem beschriebenen Fall um einen erstmaligen Vorfall und nicht um eines von vielen Beispielen handelt. Zusätzlich wird in einer der Dolmetschungen an dieser Stelle ebenfalls ausgelassen, dass „rechtliche Schritte“ (T8: Z54) vorgenommen wurden, wie es auch im Original an folgender Stelle heißt: „il a ensuite décidé d'attaquer en justice“²⁸ (T1: Z62). Diese Auslassung führt folglich zu einer zusätzlichen Sinnverfälschung, da die Tragweite dieses Falles nicht klar transferiert werden konnte.

Insgesamt finden sich in den Leistungen der Humandolmetscher:innen einige Sinnverfälschungen, jedoch beschränken sich die meisten davon auf Kleinigkeiten, die entweder für die Rezipient:innen aufgrund des Kontextes trotzdem verständlich sein sollten, oder den Sinn nur leicht verändern. Lediglich eine Stelle wird von zwei Dolmetscher:innen konträr zum Ausgangstext wiedergegeben, sodass der Sinn nicht korrekt übertragen werden konnte.

²⁶ In China zählt man ungefähr 850 Millionen Mobiltelefonnutzer.

²⁷ Diese Art von Fall ist völlig neu und ohne Präzedenz in China.

²⁸ Er hat anschließend beschlossen, vor Gericht zu gehen.

5.5.1.3 Umformulierungen

Außerdem finden sich einige lexikalische Umformulierungen, die andere Konnotationen als jene des Ausgangstextes mit sich bringen, in den Dolmetschungen, wodurch es zu leichten Sinnveränderungen in den Zieltexten kommen kann. So heißt es in einer Dolmetschung beispielsweise „weil China ein großer Fan von Technologie ist“ (T5: Z12-13), während diese Stelle im französischen Original „la Chine est souvent très friande des avancées technologiques.“²⁹ (T1: Z14-15) lautet. Demnach wird das Land China im Zieltext als ‚Fan‘ betitelt, wobei die deutsche Übersetzung hier andere Konnotationen aufweist, zu umgangssprachlich wirkt und als Attribut zu einem Land fehlplatziert wirkt. Dies wird ebenso durch die anderen Dolmetschungen bestätigt, in denen es etwa heißt, dass „China [...] versessen auf neue Technologien [ist.]“ (T6: Z12; T8: Z14). Ähnliches gilt auch für die Textstelle „problematische Bürger“ (T7: Z27) in einer Dolmetschung, während im Original von „citoyen désobéissants“ (T1: Z27) die Rede ist. Hier wird zwar ebenfalls der Inhalt transferiert, jedoch handelt es sich bei der Wortwahl der Dolmetschung um eine eher ungewöhnliche Formulierung, sodass an dieser Stelle eine wörtlichere Dolmetschung wohl passender gewesen wäre.

Weitere Umformulierungen finden sich etwa in der Zeile 27 einer Dolmetschung (T5), wo es heißt: „ein Vorschlag der Regierung“, wohingegen es im Original noch „une mesure prise par le gouvernement“³⁰ heißt, wodurch im Deutschen eventuell unklar sein könnte, ob es sich nur um eine Idee der Regierung handelt oder diese tatsächlich auch umgesetzt wurde, wie durch die Verwendung des Wortes ‚Maßnahme‘ deutlicher wird. Ähnliches ist auch in einer weiteren Dolmetschung zu erkennen, wo geäußert wird, dass „Daten einfach leaked worden sind“ (T7: Z43-44). Im Ausgangstext heißt es jedoch: „[...] qui revendent les données“³¹ (T1: Z47), sodass hier die Diskrepanz entsteht, dass im Ausgangstext die Daten willentlich verkauft wurden, während es im Zieltext den Anschein hat, dass die Daten unbeabsichtigt verbreitet wurden. Demnach ist zwar klar, dass die Daten nicht sicher sind, jedoch fehlt die klare Intention des Originals hier in der Dolmetschung, wodurch sich zusätzlich zu der Umformulierung auch ein leichter ‚information loss‘ ergibt. Im Gesamtkontext der Dolmetschung ist diese Textstelle jedoch nur als Umformulierung zu werten, da zuvor bereits an anderer Stelle erwähnt wurde, dass „Daten

²⁹ China ist oft sehr begierig nach fortgeschrittenen Technologien.

³⁰ eine von der Regierung getroffene Maßnahme

³¹ die Daten weiterverkaufen

[...] verkauft werden“ (T7: Z41-42), da es sich im Zieltext um eine redundante Information handelt.

Zusätzlich findet sich die Dolmetschung „China hat den Ruf, führend auf dem Technologiemarkt zu sein“ (T5: Z54) für „le gouvernement chinois est désireux de faire de la Chine un leader technologique“³² (T1: Z68-69). Auch hier wird durch die Umformulierung die Aussage der Textstelle leicht verändert, da der Wille der chinesischen Regierung nach einer Vorreiterrolle im Zieltext nur als Außenwahrnehmung durch die Beschreibung als ‚Ruf Chinas‘ beschrieben wird. Der grundlegende Inhalt wird folglich übermittelt, jedoch finden sich nicht alle inhaltlichen Nuancen im Zieltext wieder. Ähnliches gilt auch für die folgende Formulierung: „die Regierung [hat] geäußert, dass [...]“ (T6: Z63) als Dolmetschung von „le gouvernement a publié des directives“³³ (T1: Z75-76). Auch hier geht etwa jener inhaltliche Teilaspekt verloren, dass es sich um die Veröffentlichung von Richtlinien und nicht um eine einfache Äußerung der Regierung handelt, wobei dennoch der grundlegende Inhalt des Ausgangstextes auch im Zieltext enthalten ist.

Weitere Umformulierungen, die mit leichten Änderungen des Inhaltes einhergehen, finden sich mit „in die Privatsphäre eingreifend“ (T8: Z20-21) als Dolmetschung für „intrusive“ (T1: Z21) [aufdringlich] und „Leider gibt es keine Präzedenzfälle“ (T8: Z55) für „ce genre de cas est tout à fait nouveau“³⁴ (T1: Z63). Bei der ersten Textstelle wurde hinzugefügt, dass es sich um ein Eingreifen in die Privatsphäre handelt, sodass die Aussage im Vergleich zum Original etwas spezifiziert wurde, der Sinn aber erhalten bleibt. Beim zweiten Ausschnitt wurde seitens der Dolmetscherin durch das Hinzufügen von ‚leider‘ eine Wertung eingebaut, die so im Original nicht zu finden ist, was allerdings ebenfalls den Inhalt an diese Stelle nicht grob verändert, da es mit dem Kontext kohärent ist.

Zu all diesen eben erläuterten Umformulierungen in den Dolmetschungen sei allerdings auch erwähnt, dass es sich hierbei nicht um Sinnverfälschungen handelt und somit auch der ursprüngliche Inhalt des Ausgangstextes in die Zielsprache transferiert werden kann. So finden sich zwar Formulierungen wieder, die teilweise andere Konnotationen hervorrufen oder anderen stilistischen Ebenen zuzuordnen sind, jedoch keine groben inhaltlichen Fehler. Solche Sinnverfälschungen unterschiedlicher Schweregrade finden sich jedoch an anderen Stellen in den Dolmetschungen wieder und werden im Folgenden genauer erläutert.

³² die chinesische Regierung ist gewillt aus China einen technologischen Leader zu machen

³³ die Regierung hat Richtlinien veröffentlicht

³⁴ diese Art von Fall ist völlig neu

5.5.1.4 Kürzungen

Von allen vier Dolmetscher:innen wurden Redundanzen des Ausgangstextes in den Dolmetschungen gekürzt, um die Erwartungen an eine Konsekutivdolmetschung, die kürzer als das Original sein sollte, erfüllen zu können. Ein erstes Beispiel für die Kürzung von Redundanzen findet sich etwa bei der Stelle „Doch es gibt sehr viele Vorteile.“ (T5: Z20), in der im Original zusätzlich noch genauer ausgeführt wird, dass einer der erwähnten Vorteile das Gewährleisten der öffentlichen Sicherheit darstellt:

Selon eux, la reconnaissance faciale leur facilite la vie et elle est synonyme de nombreux bénéfices. En plus, selon eux, la reconnaissance faciale permet d'assurer la sécurité publique.³⁵ (T1: Z22-23)

In der Originalrede wird dieses Detail allerdings im Folgesatz wiederholt, wo es sich auch in der Dolmetschung wiederfindet, sodass hier kein ‚information loss‘ stattfindet. Ähnliches gilt beispielsweise auch für den Abschnitt des Ausgangstextes, in dem es um die Gesichtserkennung bei Nutzung von Mobiltelefonen geht (vgl. T1: Z31-38), in welchem drei Dolmetscher:innen (vgl. T6: Z29-34; T7: Z31-36; T8: 30-34), die vorkommenden Redundanzen auf eine solche Art und Weise gekürzt haben, dass es zu keinen inhaltlichen Einbußen kommt.

5.5.1.5 Kohärenz, Kohäsion, Pragmatik und Terminologie

In Bezug auf die Kohärenz lässt sich festhalten, dass alle vier Dolmetschleistungen in sich kohärent und logisch stimmig sind. Somit kommen innerhalb der einzelnen Zieltextes keine Widersprüche und auch keine falschen Verknüpfungen zwischen einzelnen Satzteilen vor, die inhaltliche Verzerrungen mit sich bringen, wie es bei Google Übersetzer der Fall war. Außerdem wurden die textuellen Verknüpfungen adäquat gewählt und entsprechen hinsichtlich des Zwecks auch jenen des Ausgangstextes, sodass auch die Kohäsion der einzelnen Dolmetschleistungen gegeben ist. Hierbei lässt sich folglich bereits ein starker Unterschied zu der maschinellen Dolmetschung feststellen, bei der es – wie im vorangegangenen Kapitel ausführlicher erläutert wurde – zu einigen Widersprüchen aufgrund fehlerhafter Spracherkennung kam. Bei den menschlichen Dolmetscher:innen hingegen hilft Kontextwissen, um sich logische

³⁵ Laut ihnen erleichtert die Gesichtserkennung ihr Leben und bringt viele Vorteile mit sich. Außerdem erlaubt die Gesichtserkennung laut ihnen, die öffentliche Sicherheit zu gewährleisten.

Zusammenhänge erschließen zu können, auch wenn vielleicht nicht der gesamte Inhalt des Ausgangstextes lückenlos verstanden werden konnte oder im Laufe des Konsektivdolmetschprozesses in Erinnerung geblieben ist.

Dieses Kontextwissen geht ebenso Hand in Hand mit situativem Wissen, was für pragmatische Anpassungen im Zieltext benötigt wird, um eine akzeptable Dolmetschung produzieren zu können. Hierbei lassen sich allerdings Unterschiede zwischen den einzelnen Dolmetschungen feststellen. Dies ist vor allem bei der Frage nach dem Gendern im Deutschen auffallend. So wird von einer Dolmetscherin (T5) bis auf eine einzige Ausnahme im gesamten Zieltext entweder mit einer Doppelnennung oder mit der Genderform mit Glottisschlag gegendert. Eine weitere Dolmetscherin (T8) nutzt bis auf zwei Ausnahmen genderneutrale Umschreibungen, um das generische Maskulinum zu vermeiden, während die anderen beiden (T6, T7) nicht gendern und die Formen des generischen Maskulinums wie im Ausgangstext und bei der maschinellen Dolmetschung verwenden. Dementsprechend ist in Bezug auf pragmatische Anpassungen an die Zielkultur ein Ungleichgewicht festzustellen, da sich nicht alle Dolmetscher:innen für die gleiche Vorgehensweise entschieden haben. Insgesamt ist dennoch ein Unterschied zur maschinellen Dolmetschung festzustellen, bei der eine solche Anpassung an geschlechtergerechte Sprache in keiner Weise erfolgte.

Ein weiterer Punkt bezüglich der Pragmatik betrifft die angemessene Begrüßung und Verabschiedung, wobei hier festzuhalten ist, dass eine Dolmetscherin (T6) sowohl die Begrüßung als auch die Verabschiedung ausließ, während eine weitere (T8) auf die Verabschiedung verzichtete. Die anderen beiden allerdings haben adäquate Begrüßungs- sowie Verabschiedungsfloskeln gewählt. Bei der maschinellen Dolmetschung von Google Übersetzer hingegen ist die Begrüßung – wie oben erläutert – als etwas informell einzustufen. Außerdem ist bei den Humandolmetscher:innen positiv hervorzuheben, dass alle eine Umschreibung für das Wort ‚Führer‘ verwendeten, das im Zieltext der maschinellen Dolmetschung vorkam. Demnach kam es auch hier zu einer passenden pragmatischen Anpassung seitens der Dolmetscher:innen, um dieses im Deutschen stark negativ konnotierte Wort zu umgehen.

Bezüglich der korrekten Verwendung von Terminologie gilt klarerweise ebenfalls, dass der vorliegende Ausgangstext keine besonderen terminologischen Schwierigkeiten in sich birgt. Demnach sollte lediglich der Terminus ‚Gesichtserkennung‘ einheitlich verwendet werden, was bei allen Dolmetschungen – abgesehen von jeweils einer Ausnahme bei zwei Dolmetschungen (T7: Z58; T8: Z56), wo „Gesichtererkennung“ beziehungsweise „Gesichtstechnologien“ verwendet wurde – der Fall war, wobei dies auch bei der maschinellen Dolmetschung kein Problem darstellte.

5.5.2 Realisierung der Dolmetschung

5.5.2.1 Stimme, Sprechtempo und Rhythmus

In Bezug auf die Realisierung der Dolmetschung ist grundsätzlich der Unterschied zwischen der synthetischen und hinsichtlich des Klangs stark von natürlichen Stimmen abweichenden Stimme der maschinellen Dolmetschung und den menschlichen Stimmen der Dolmetscher:innen hervorzuheben. Beim Vergleich zwischen den einzelnen Dolmetscher:innen lassen sich aber auch hier Unterschiede ausmachen. So sprechen zwar alle vier mit einer angenehm zu hörenden und lebhaften Stimme, jedoch variiert die Anzahl an Füllwörtern und Verzögerungslaute je nach Dolmetschung. Bei der ersten Dolmetscherin (T5) sind kaum Füllwörter und Verzögerungslaute auszumachen, während bei der zweiten (T6) viele Verzögerungslaute und auch einige Selbstkorrekturen, sowie weitere Hintergrundgeräusche, wie etwa das Rascheln von Papier, zu hören sind, was als Störfaktoren für die Rezipient:innen gelten und das Zuhören erschweren kann. Bei der dritten und vierten Dolmetschung (T7; T8) sind wiederum nur wenige Füllwörter und lediglich einige Verzögerungslaute zu erkennen. Insgesamt sind die Stimmen der Dolmetscher:innen dementsprechend vor allem aufgrund des Klangs und auch wegen der wenigen Verzögerungslaute und Füllwörter wohl als positiver als der künstliche Klang der Sprachausgabe von Google Übersetzer zu bewerten.

Beim Sprechtempo fällt auf, dass dieses bei allen vier Dolmetschleistungen – wie auch bei der maschinellen Dolmetschung – zumindest etwas langsamer als das französischsprachige Original ist. Bei einer Dolmetschung (T6) ist darüber hinaus anzumerken, dass zusätzlich zum langsamen Sprechen auch einige Verzögerungen durch einige längere Pausen von ein bis zwei Sekunden, die Artikulation von Verzögerungslaute und das Auftreten von außersprachlichen Störgeräuschen, wie häufiges Räuspern oder Husten, wodurch sich die Dolmetschung zusätzlich verlängert. Bei einer weiteren (T8) ist das Tempo etwas zu langsam, wodurch sich trotz der Auslassung von Redundanzen im Ausgangstext auch hier die insgesamt Länge der Dolmetschung ausdehnt. So entspricht diese Dolmetschung mit einer Länge von 7 Minuten und 7 Sekunden etwa 114% der Länge des Ausgangstextes und ist somit für eine optimale Konsekutivdolmetschung, die – wie in Kapitel 2.3 genauer erläutert wurde – etwa 70% des Originaltextes in Anspruch nehmen sollte, als zu lange zu erachten. Ähnliches gilt auch für zwei weitere Dolmetschungen (T6 & T7), die mit einer Länge von 6 Minuten und 10 Sekunden beziehungsweise 5 Minuten und 52 Sekunden jeweils 99% und 94% der Länge des Originals dauern, was

ebenfalls als etwas zu lange zu erachten ist, um die Erwartungen der Rezipient:innen erfüllen zu können. Bei der ersten Dolmetschung hingegen (T5) wurde die angestrebte Länge durch ein angemessenes Sprechtempo und entsprechende Kürzungen von Redundanzen mit 4 Minuten und 26 Sekunden, was 71% der Länge des Ausgangstextes entspricht, sehr gut getroffen.

Hinsichtlich des Rhythmus ist festzuhalten, dass drei der vier Dolmetschungen in einer gleichmäßigen Sprechweise erfolgen, die durch wenige Verzögerungen gekennzeichnet ist. Bei der ersten Dolmetscherin (T5) fällt darüber hinaus auf, dass keine Füllwörter sowie nur recht wenige überlange Pausen vorkommen, die einen angenehmen und gleichmäßigen Rhythmus stören könnten. Ähnliches gilt für die dritte Dolmetschung (T7), wobei hier zusätzlich anzumerken ist, dass durchaus einige überlange Pausen auftreten sowie einige weitere Verzögerungen vorkommen, die durch Selbstkorrekturen oder Verzögerungslaute bedingt sind, wobei andererseits kaum Füllwörter auszumachen sind. Bei der vierten Dolmetschung (T8) lässt sich wiederum feststellen, dass einige Verzögerungen durch die langsame Sprechweise, sowie einige wenige überlange Pausen und Selbstkorrekturen vorhanden sind, wobei abermals kaum Füllwörter auftreten. Bei der zweiten Dolmetschung (T6) hingegen wirkt der Rhythmus beziehungsweise die Sprechweise insgesamt etwas abgehackt, da es doch zu stärkeren Verzögerungen durch eine Vielzahl an Verzögerungslauten, Selbstkorrekturen sowie weiteren Störgeräuschen, wie etwa Husten oder Räuspern kommt. Dafür finden sich bei dieser Dolmetschung recht wenige überlange Pausen und kaum Füllwörter.

5.5.2.2 Aussprache, Intonation und Prosodie

In Bezug auf die Aussprache der Dolmetscher:innen lässt sich feststellen, dass bei allen eine deutliche und klare Aussprache ohne Nuscheln oder andere störenden Faktoren vorhanden ist, sodass die dargebotenen Dolmetschungen gut verständlich sind. Hierbei lässt sich lediglich anmerken, dass bei einer Dolmetscherin (T7) Deutsch nicht die Erstsprache darstellt, was sich allerdings nicht störend äußern dürfte, da die Aussprache sehr passend ist und lediglich ein leichter fremdsprachlicher Akzent erkennbar ist, der das Verständnis allerdings nicht negativ beeinflusst und seitens Rezipient:innen – wie in Kapitel 2.2 beschrieben wurde – normalerweise nicht als Störfaktor wahrgenommen wird.

Auch die Intonation sowie die Prosodie der Dolmetschungen betreffend sind alle vier Leistungen ähnlich zu bewerten. So findet sich bei allen eine lebhaftes Intonation, sodass kein monotones Sprechen erkennbar ist. Weiters ist die Satzmelodie deutlich zu erkennen und auch

die Stimmhöhe wird entsprechend angepasst, sodass Fragesätze und auch die Satzgrenzen zwischen den einzelnen Aussagesätzen hervorstechen, da die Stimmhöhe entsprechend angepasst wird, sodass das mühelose Zuhören für die Rezipient:innen erleichtert wird. Lediglich bei einer Dolmetschung (T6) wirken die vielen Pausen und Verzögerungslaute, die oben genauer beschrieben wurden, auch etwas störend.

5.5.2.3 Grammatik

Bei der Analyse der Grammatik der Dolmetschungen ist festzuhalten, dass bei der ersten Dolmetschung (T5) keine Grammatikfehler vorhanden sind und die Dolmetschung stets aus vollständigen und abgeschlossenen Sätzen besteht. Bei den übrigen Dolmetschungen finden sich lediglich einige kleinere Grammatikfehler, wobei es sich abgesehen davon ebenfalls um vollständige, grammatikalisch korrekte und abgeschlossene Sätze handelt. Zu diesen Grammatikfehlern zählt etwa die falsche Verwendung von anaphorischen Pronomen, wie etwa bei der Verwendung des neutralen Pronomens „es“ (T6: Z3, Z6, Z16, Z24; T7: Z21) als Referenz auf das feminine Substantiv ‚Gesichtserkennung‘. Ein weiteres Beispiel hierfür ist „seine verbreitete Benutzung“ (T7: Z2), wobei es sich ebenfalls um eine Referenz auf die eben genannte ‚Gesichtserkennung‘ handelt, weswegen hier abermals das Genus nicht korrekt angepasst wurde.

Außerdem finden sich Beispiele für fehlerhafte Flexion, wie etwa „eine letztes Beispiel“ (T6: Z50), wo der Artikel nicht das korrekte Genus aufweist, oder „[...] dass sie [Anm. die Technologie] das Leben erleichtern [...] und sicherstellen.“, wo das Verb im falschen Numerus zum Substantiv flektiert wurde. Hinzu kommt die falsche Pluralbildung „Irisen“ (T7: Z48), sowie die Formulierung „bereitet schon für viele Sorgen“ (T7: Z29) anstelle von ‚bereitet vielen Sorge‘, bei der folglich eine nicht passende Präposition für das Verb gewählt wurde. Hinsichtlich all dieser Grammatikfehler ist aber festzuhalten, dass diese zwar wohl den Rezipient:innen auffallen und eventuell als störend empfunden werden könnten, allerdings das Verständnis des Inhalts in diesen konkret hier auftretenden Fällen nicht negativ beeinflussen sollten.

5.5.2.4 Syntax

Bei der Analyse der Syntax ist festzuhalten, dass vor allem Fehler in der Satzstellung auftauchen, die in schriftlichen Texten nicht adäquat wären, in der gesprochenen Sprache allerdings durchaus vorkommen und demnach auch als akzeptabel erachtet werden können. Den Rezipient:innen dürften diese vermeintlichen Fehler wohl nicht als störend auffallen, da auch bei der Produktion von deutschsprachigen Originaltexten solche Ausdrucksweisen vorkommen und sich dies demnach nicht auf Dolmetschungen beschränkt. Nichtsdestotrotz soll auch auf diese syntaktischen Ungenauigkeiten kurz eingegangen werden. So findet sich etwa der folgende Satz, in dem das Verb „zu kontrollieren“ vor dem anschließenden Objekt steht, wobei dieses in der Schriftsprache dahinter platziert werden müsste. Für den mündlichen Sprachgebrauch kann diese Formulierung aber wohl akzeptiert werden, da sie in jedem Fall verständlich ist:

So geht es zum Beispiel schneller, die Identität zu kontrollieren in öffentlichen Transportmitteln und Menschen, die sich nicht an die Regeln halten, zu identifizieren. (T5: Z24-26)

Ähnliches gilt für „Zum Beispiel, wenn man sich ein neues Handy kauft, [...] soll das Gesicht verpflichtend registriert werden“ (T6: Z30-32), wo die Satzstellung ebenfalls der mündlichen Umgangssprache entlehnt wurde. In der gesprochenen Sprache können ebenso Einschübe akzeptiert werden, auch wenn diese in der Schriftsprache syntaktisch nicht als korrekt erachtet werden. Ein weiteres solches Beispiel findet sich etwa im Satz „Heutzutage muss man seine Gesichtszüge registrieren und auch die Telefonnummer.“ (T8: Z31-32), in dem am Ende noch ein Zusatz erfolgt, der gemäß einer korrekten Satzstellung vor dem zweiten Teil des Prädikats zu erfolgen hat, wobei auch dies im mündlichen Gebrauch nicht ungewöhnlich ist und somit akzeptiert werden kann.

Hinzu kommt, dass auch einige Sätze in den Dolmetschungen vorkommen, die abgesehen von mündlich wohl akzeptierten Formulierungen weitere Fehler enthalten, die den Rezipient:innen auffallen könnten und dadurch die Erwartungshaltung an eine sprachlich korrekte Dolmetschung stören könnten, auch wenn der Inhalt dabei korrekt übertragen wird. Dazu zählt etwa der folgende Satzausschnitt:

Und eine letztes Beispiel, welches ebenfalls das Vertrauen der chinesischen Bevölkerung sehr senkt, ja sehr verringert hat, ist das Beispiel, wo ein Professor in ein Wildtierreservat in der Nähe von Shanghai sich geweigert hat, sein Gesicht scannen zu lassen und er ist deswegen nicht in das Reservat eingetreten und ist dann juristisch gegen das Reservat vorgegangen und [...]. (T6: Z52-54)

Dieser stellt ein Beispiel für einen solchen syntaktisch nicht vollends korrekten Satz dar, da etwa die Prädikate nicht an den richtigen Positionen im Satz platziert wurden, wodurch der Inhalt zwar korrekt wiedergegeben wurde, es allerdings beim Zuhören etwas verwirrend wirken könnte, da die Erwartungen an die Satzstellung nicht erfüllt werden. Ähnliches gilt auch für die beiden Sätze „[...] denn China ganz viele Fortschritte [...] macht.“ (T7: Z14-15) und „[...] jedes Mal, wenn jemand ein Handy kauft durch einen Gesichtserkennungsprozess zu gehen.“ (T7: Z32), in welchen ebenso das Prädikat nicht an der richtigen Stelle im Satz positioniert wurde. Weitere syntaktische Fehler finden sich in einem Satz mit einer doppelten Nennung des Prädikats „möchte“ (T7: Z33-34), sowie mit der ungrammatischen Formulierung des Passivs „dafür verdächtigt werden“ (T7: Z49) und der fehlerhaften Verwendung der Präposition ‚seit‘ in „seit 2013 in Kraft getreten ist“ (T7: Z35-36). Bei diesen Fehlern handelt es sich aber ebenfalls um solche, die den Rezipient:innen zwar negativ auffallen beziehungsweise nicht deren Erwartungen entsprechen könnten, die aber gleichzeitig das Verständnis des gedolmetschten Inhaltes insgesamt nicht negativ beeinflussen sollten.

5.5.2.5 Lexik und Stil

Bezüglich der Bewertung der Lexik der Dolmetschleistungen der Humandolmetscher:innen fallen sowohl positive als auch negative Punkte bei der Analyse auf. Positiv zu bewerten sind etwa Formulierungen, bei denen es sich nicht um wörtliche Übertragungen von Äußerungen des Ausgangstextes handelt, sondern die vielmehr passend für die Zielsprache und Zielkultur abgeändert wurden. Dazu zählt beispielsweise die Übersetzung „führend“ (T5: Z54) für das originalsprachliche „leader“ (T2: Z62) beziehungsweise die direkte Übernahme des Anglizismus „Leader“ (T8: Z60) auch im deutschen Zieltext. Eine wörtliche Übersetzung mit dem Wort „Führer“ (T3: Z62), wie dies bei der maschinellen Übersetzung der Fall war, ist – wie bereits erläutert – im Deutschen aufgrund der starken negativen Konnotation nicht als adäquat zu erachten. Weiters positiv hervorzuheben sind die Formulierungen „in einem positiven Licht“ (T7: Z20) als Übersetzung für „[ils] semblent assez favorables“ (T1: Z20), „als Referenz anführen“ (T8: Z13) für „faire référence à“ (T1: Z13-14) und „rechtliche Schritte vornehmen“ (T8: Z54) für „attaquer en justice“ (T1: Z62), da auch in diesen Beispielen der Sinn des Ausgangstextes wiedergegeben wurde und gleichzeitig eine passende Formulierung losgelöst von jener des Originals gewählt wurde.

Andererseits finden sich in den Dolmetschung ebenso einige Formulierungen, die auf der lexikalischen Ebene nicht als vollends korrekt zu erachten sind. So findet sich zum Beispiel die Formulierung „Platz in unserem Leben zu nehmen“ (T7: Z6), wobei hier die Verwendung des Verbs ‚einzunehmen‘ adäquater wäre. Ähnliches gilt für „Wer sich [...] wünscht, [...] ein Passwort einzugeben“ (T7: Z10-11), wo beispielsweise die Verwendung von ‚mögen‘ oder ‚wollen‘ anstelle von ‚wünschen‘ passender gewesen wäre. Eine weitere eher unpassende Formulierung lautet „durch einen Gesichtserkennungsprozess zu gehen“ (T7: Z32-33) anstelle von beispielsweise „[sie müssen sich] mit ihrem Gesicht registrieren“ (T5: Z32-33) als Dolmetschung für „faire enregistrer leurs visages“ (T1: Z 34). Ähnliches gilt für „Daten [...] werden [...] gemacht“ (T7: Z47-48), wo abermals anstelle von etwa „Daten sammeln“ (T5: Z46) ein eher inadäquates Verb verwendet wurde. Bei all diesen Beispielen gilt aber wiederum, dass diese zwar lexikalisch nicht ganz korrekt für den Zieltext sind, nichtsdestotrotz dadurch keine Verständnisprobleme bei den Rezipient:innen entstehen sollten, weswegen der Inhalt des Ausgangstextes trotzdem in die Zieltext transportiert werden kann.

Hinsichtlich des Stils der Dolmetschungen lässt sich insgesamt festhalten, dass alle vier Dolmetschungen dem Ausgangstext bezüglich des Formalitätsgrades entsprechen, weswegen der verwendete Stil jeweils als adäquat für den Zweck und den Kontext der Dolmetschung zu erachten ist. Eine kleine Ausnahme davon stellt lediglich die Aussprache des Ausdruckes „Aug“ (T6: Z46) dar, da diese aufgrund des fehlenden Schwa-Lautes am Wortende recht umgangssprachlich erscheint und demnach nicht dem erfordernten Stilniveau entspricht. Bei gesamter Betrachtung aller Dolmetschleistungen fällt dieser einzelne Punkt allerdings nur sehr schwach ins Gewicht. Auch hier ist anzumerken, dass der französischsprachige Ausgangstext keine stilistischen Besonderheiten enthält, auf die bei einer Dolmetschung besonders geachtet werden muss oder die im Rahmen der Erstellung eines adäquaten Zieltextes angepasst werden müssten.

5.5.3 Quantitative Analyse

In diesem Kapitel folgt nun die quantitative Analyse der Fehler der menschlichen Dolmetscher:innen. Diese sollen zunächst einzeln pro Dolmetschung aufgelistet werden, ehe anschließend ein arithmetisches Mittel errechnet werden soll, damit dieses schlussendlich mit jenem von Google Übersetzer verglichen werden kann.

Hauptkategorie	Fehlerkategorie	Leicht	Mittel	Schwer	Gesamt
Übertragung des Inhalts	Treue zum Ausgangstext	6	3	1	10
Realisierung der Dolmetschung	Syntax	1	0	0	1
	Lexik	1	1	0	2
Gesamt		8	4	1	13

Fehleranalyse Dolmetschung 1

Bei Dolmetschung 1 (T5) finden sich – wie in der obenstehenden Tabelle sichtbar ist – bezüglich der Übertragung des Inhalts sechs leichte Fehler, drei mittlere Fehler und ein schwerer Fehler, die sich allesamt auf die Treue zum Ausgangstext beziehen, da keine Fehler hinsichtlich der Kohärenz und der Pragmatik erkennbar sind. Die leichten Fehler (vgl. T5: Z47, Z53, Z62, Z62-63, Z69-70, Z78) beruhen dabei jeweils auf der Auslassung von Details beispielsweise in Aufzählungen, die den Rezipient:innen demnach zwar etwas Inhalt vorenthalten, das Gesamtverständnis des Inhaltes allerdings nicht stark negativ beeinflussen. Bei den mittleren Fehlern handelt es sich um leichte Sinnveränderungen (vgl. T5: Z46, Z54) beziehungsweise um eine Auslassung (vgl. T5: Z73) mit stärkerem negativen Einfluss auf die Verständlichkeit als bei den oben genannten. Der schwere Fehler beruht auf einer Sinnumkehr (vgl. T5: Z51), in Folge derer der Inhalt an dieser Stelle für die Rezipient:innen nicht mehr vollends verständlich ist.

In Bezug auf die Realisierung der Dolmetschung 1 (T5) finden sich zwei leichte und ein mittlerer Fehler, wobei sich keine Fehler hinsichtlich der Grammatik feststellen lassen. So finden sich ein leichter Fehler bezüglich der Satzstellung (vgl. T5: Z25), sowie ein leichter (vgl. T5: Z13) und ein mittlerer Fehler (vgl. T5: Z27) bezüglich der Lexik, da an diesen Stellen Wörter mit leicht beziehungsweise stärker ausgeprägter anderer Konnotation als im Original verwendet wurden. Insgesamt ergeben sich somit bei Dolmetschung 1 acht leichte, vier mittlere und ein schwerer Fehler.

Hauptkategorie	Fehlerkategorie	Leicht	Mittel	Schwer	Gesamt
Übertragung des Inhalts	Treue zum Ausgangstext	6	0	0	6
	Pragmatik	0	2	0	2
Realisierung der Dolmetschung	Grammatik	1	4	0	5
	Syntax	1	1	0	2
	Lexik	3	1	0	4
Gesamt		11	8	0	19

Fehleranalyse Dolmetschung 2

Bei Dolmetschung 2 (T6), deren tabellarische Auflistung der Fehler oben zu sehen ist, finden sich in Bezug auf die Übertragung des Inhalts sechs leichte und zwei mittlere Fehler, wobei sich keine bei der Kohärenz feststellen lassen. Die sechs leichten Fehler können alle der Treue zum Ausgangstext zugeordnet werden, wobei es sich in fünf Fällen (vgl. T6: Z9-12, Z46-48, Z53, Z63, Z69-71) um Auslassungen von Details handelt, die keine allzu großen Auswirkungen auf die Verständlichkeit und die korrekte Übertragung des Inhaltes haben. Bei einem weiteren leichten Fehler (vgl. T6: Z24) handelt es sich um eine Spezifizierung, deren Auswirkungen für die Nutzer:innen allerdings ebenfalls als minimal zu erachten sind. In Bezug auf die Pragmatik lässt sich festhalten, dass sowohl die Begrüßung als auch die Verabschiedung fehlt, was jeweils als mittlerer Fehler gewertet wurde.

In Bezug auf die Grammatik lassen sich ein leichter Fehler (vgl. T6: Z50) aufgrund der falschen Flexion eines Artikels sowie vier mittlere Fehler (vgl. T6: Z1, Z3, Z6, Z24; Z50) aufgrund der Verwendung des falschen Pronomens als Referenz auf ein zuvor genanntes Substantiv festhalten. Weiters finden sich ein leichter (vgl. T6: Z30) und ein mittlerer Fehler (vgl. T6: Z52-54) in Bezug auf die Syntax, die auf falschen Satzstellungen beruhen, wobei der mittlere Fehler das Verständnis bei den Rezipient:innen stärker negativ beeinflusst, als dies bei dem leichten Fehler der Fall ist. Auch bei der Lexik finden sich drei leichte (vgl. T6: Z5, Z12, Z18) und ein mittlerer Fehler (vgl. T6: Z37), die auf der Verwendung von Wörtern mit leicht beziehungsweise etwas stärkerer anderer Konnotation als im Original beruhen und demnach das Verständnis des Inhaltes für die Rezipient:innen leicht oder mittelstark negativ beeinflussen

können. Insgesamt ergeben sich somit bei der Dolmetschung 2 (T6) elf leichte Fehler und acht mittlere Fehler.

Hauptkategorie	Fehlerkategorie	Leicht	Mittel	Schwer	Gesamt
Übertragung des Inhalts	Treue zum Ausgangstext	5	5	0	10
Realisierung der Dolmetschung	Grammatik	3	2	0	5
	Syntax	5	0	0	5
	Lexik	2	6	0	8
Gesamt		15	13	0	28

Fehleranalyse Dolmetschung 3

Bei Dolmetschung 3 (T7) finden sich – wie in der obigen Tabelle erkennbar ist – bezüglich der Übertragung des Inhaltes fünf leichte und fünf mittlere Fehler. Hinsichtlich der Treue zum Ausgangstext finden sich demnach fünf leichte Fehler (vgl. T7: Z25-29, Z42-43, Z53, Z64, Z76), die allesamt auf dem Auslassen von Details beruhen. Drei der mittleren Fehler (vgl. T7: Z19, Z23, Z43-44) basieren auf leichten Sinnverfälschungen, die mit entsprechenden Einbußen bezüglich der Qualität der Dolmetschung einhergehen. Bei den übrigen zwei mittleren Fehlern handelt es sich zum einen um eine etwas unklar formulierte Stelle (vgl. T7: Z18) und zum anderen um eine Verallgemeinerung (vgl. T7: Z69-71), weswegen hier ebenso inhaltliche Einbußen zu erkennen sind.

In Bezug auf die Realisierung der Dolmetschung 3 (T7) lassen sich zehn leichte und acht mittlere Fehler erkennen. Bezüglich der Grammatik finden sich drei leichte und zwei mittlere Fehler. Zwei der leichten Fehler beruhen auf der Verwendung einer falschen Präposition (vgl. T7: Z29, Z35-36), der dritte auf einer falschen Pluralform (T7: Z48). Da die betreffenden Textpassagen trotz der vorkommenden Fehler inhaltlich gut verständlich sind und demnach für die Rezipient:innen nicht allzu störend wirken, wurden diese Fehler als leicht klassifiziert. Bei den beiden mittleren Fehler handelt es sich um die Verwendung von falschen Pronomen als Referenz auf zuvor genannte Substantive (vgl. T7: Z2, Z21), was zu Verständnisschwierigkeiten führen könnte. In Bezug auf die Syntax finden sich fünf leichte Fehler, wovon vier auf

Fehler im Satzbau fallen (vgl. T7: Z14-15, Z32, Z49, Z55), während sich auch eine doppelte Nennung eines Verbes erkennen lässt (vgl. T7: Z33-34), was als weiterer leichter Fehler kategorisiert wurde. Zwei leichte und sechs mittlere Fehler betreffen die Lexik der Dolmetschung. Ein leichter Fehler basiert auf der Verwendung einer unpassenden Präposition (vgl. T7: Z6). Der zweite leichte Fehler (vgl. T7: Z11), sowie die sechs mittleren Fehler (vgl. T7: Z27, Z32-33, Z47-48, Z53, Z54, Z56) beruhen auf unpassenden Formulierungen, wobei die mittleren Fehler größere negative Auswirkungen auf die Qualität der Dolmetschleistung haben, weswegen die Fehler hier in die zwei Kategorien aufgeteilt wurden. Insgesamt finden sich somit bei der Dolmetschung 3 (T7) fünfzehn leichte und dreizehn mittlere Fehler.

Hauptkategorie	Fehlerkategorie	Leicht	Mittel	Schwer	Gesamt
Übertragung des Inhalts	Treue zum Ausgangstext	2	3	1	6
	Pragmatik	0	1	0	1
Realisierung der Dolmetschung	Grammatik	2	0	0	2
	Syntax	1	0	0	1
Gesamt		5	4	1	10

Fehleranalyse Dolmetschung 4

Bei Dolmetschung 4 (T8) lassen sich – wie in der obigen Tabelle sichtbar wird – hinsichtlich der Übertragung des Inhalts insgesamt zwei leichte Fehler, vier mittlere Fehler und ein schwerer Fehler ausmachen. Dabei finden sich keine Fehler bezüglich der Kohärenz. Ein pragmatischer Fehler betrifft die fehlende Verabschiedung (vgl. T8: Z70), was als mittlerer Fehler gewertet wird. Bei Betrachtung der Treue zum Ausgangstext fallen zwei leichte Fehler auf, wobei einer davon leichte Hinzufügungen von Informationen (vgl. T8: Z20-21) und der andere eine leichte Sinnveränderung (vgl. T8: Z56-57) betrifft. Zwei der mittleren Fehler basieren auf Sinnveränderungen mit größerer negativer Auswirkung (vgl. T8: Z44, Z55) als jener leichte Fehler. Der dritte mittlere Fehler beruht auf einer Auslassung von Informationen (vgl. T8: Z73) des Ausgangstextes. Beim schweren Fehler handelt es sich um eine schwerwiegende Sinnveränderung

(vgl. T8: Z46), da an dieser Stelle eine Fehlübersetzung passiert, sodass für die Rezipient:innen ein anderer Inhalt als im Ausgangstext transferiert wird.

In Hinblick auf die Realisierung der Dolmetschung 4 (T8) lassen sich drei leichte Fehler erkennen. Bei der Grammatik finden sich zwei leichte Fehler, die auf einer jeweils fehlerhaften Flexion eines Verbes (vgl. T8: Z21, Z22) basieren. Ein weiterer leichter Fehler betrifft den Satzbau und damit die Syntax (vgl. T8: Z21-22). Durch diese drei Fehler wird jedoch das Verständnis für die Rezipient:innen nicht stark negativ beeinflusst, weswegen hier nur leichte Fehler gezählt werden. Bezüglich der Lexik wiederum finden sich keine Fehler. Insgesamt lassen sich somit bei der Dolmetschung 4 fünf leichte Fehler, vier mittlere Fehler und ein schwerer Fehler erkennen.

Hauptkategorie	Fehlerkategorie	Leicht	Mittel	Schwer	Gesamt
Übertragung des Inhalts	Treue zum Ausgangstext	4,75	2,75	0,5	8
	Pragmatik	0	0,75	0	0,75
Realisierung der Dolmetschung	Grammatik	1,5	1,5	0	3
	Syntax	2	0,25	0	2,25
	Lexik	1,5	2	0	3,5
Gesamt		9,75	7,25	0,5	17,5

Arithmetisches Mittel aus den Fehleranalyse der vier menschlichen Dolmetschungen

Basierend auf der hier durchgeführten quantitativen Fehleranalyse der menschlichen Dolmetschleistungen ergeben sich folgende arithmetische Mittelwerte, die in der obigen Tabelle zusammengefasst sind und die in weiterer Folge neben den Ergebnissen der qualitativen Analyse für den Vergleich zwischen Mensch und Maschine herangezogen werden sollen: In Bezug auf die Übertragung des Inhaltes finden sich durchschnittlich 4,75 leichte, 2,75 mittlere und 0,5 schwere Fehler in der Kategorie der Treue zum Ausgangstext und 0,75 mittlere Fehler bei der Pragmatik. In Bezug auf die Realisierung der Dolmetschung finden sich in der Fehlerkategorie Grammatik durchschnittlich 1,5 leichte und 1,5 mittlere Fehler. In Bezug auf die Syntax finden sich durchschnittlich zwei leichte und 0,25 mittlere Fehler und hinsichtlich der Lexik lassen

sich im Durchschnitt 1,5 leichte und zwei mittlere Fehler erkennen, wobei in keiner Dolmetschung ein schwerer Fehler bei der Realisierung der Dolmetschung auszumachen ist.

5.5.4 Gesamtbewertung

Zusammenfassend lässt sich hinsichtlich der Bewertung der Leistungen der Humandolmetscher:innen festhalten, dass – wie bereits vermutet – einige Unterschiede in den verschiedenen Ebenen der Bewertung der Qualität der Dolmetschleistungen auszumachen sind. Diese Annahme stellte – wie in Kapitel 4 erläutert wurde – die Grundlage dafür da, nicht nur eine, sondern mehrere Humandolmetschungen zur Analyse heranzuziehen, um herausstechende Einzelfälle gesondert betrachten zu können und einen generellen Überblick über die Leistung von menschlichen Dolmetscher:innen zu erhalten, ohne dass einzelne Fehler zu stark ins Gewicht fallen.

Unter Berücksichtigung dessen lässt sich in Bezug auf die Bewertung dieser Dolmetschungen festhalten, dass bezüglich der Treue zum Ausgangstext alle produzierten Zieltexte die Wirkung des Ausgangstextes auf eine solche Art und Weise transferieren konnten, dass der grundlegende Zweck der Dolmetschung erfüllt ist, da die relevanten Informationen in allen Dolmetschungen vorhanden sind und es zu keinen wirklichen Widersprüchen zwischen dem Ausgangstext und den Zieltexten kommt. Nichtsdestotrotz finden sich auch einige kleinere Unstimmigkeiten wieder, die vor allem den Informationsverlust betreffen. So wurden insgesamt einige Details weggelassen oder auch Formulierungen etwas allgemeiner gehalten, als dies im Ausgangstext der Fall war.

Alle Dolmetschungen sind außerdem in sich kohärent, logisch stimmig, widerspruchsfrei und durch adäquate Verknüpfungen entsprechend des Ausgangstextes strukturiert. Bezüglich des Genders finden sich unterschiedliche Herangehensweisen, sodass teilweise gegendert wird, wohingegen in anderen Dolmetschungen das generische Maskulinum wie im Ausgangstext verwendet wird und somit keine pragmatische Anpassung an das Zielpublikum erfolgt. Der Ausgangstext ist hinsichtlich der verwendeten Terminologie recht einfach gehalten, weswegen dieser Punkt auch in den Dolmetschungen keine Probleme bereitet und die verwendeten Termini korrekt und konsequent eingesetzt werden.

Bezüglich der Realisierung der Dolmetschung ist festzuhalten, dass alle Dolmetscher:innen über eine angenehme und lebhafte Stimme verfügen und großteils nur wenige

Füllwörter, Verzögerungslaute und Störgeräusche vorkommen. Allerdings ist das Sprechtempo teilweise etwas zu langsam, wodurch sich auch die insgesamt zu langen Dolmetschungen ergeben. Dafür sind die Sprechweisen durch einen gleichmäßigen Rhythmus geprägt und die Aussprache ist klar und gut verständlich. Auch bezüglich der Prosodie und Intonation werden die Erwartungen an eine deutlich erkennbare Satzmelodie erfüllt. Grammatikfehler sind ebenso kaum vorhanden und beschränken sich nahezu ausschließlich auf die seltene Verwendung von falschen Pronomen als Referenz auf zuvor verwendete Substantive. Die Dolmetschungen bestehen sonst aus vollständigen und abgeschlossenen Sätzen, die lediglich manchmal durch die fehlerhafte Stellung des Prädikats gestört werden, wobei es sich dabei meist um Formulierungen handelt, die im mündlichen Sprachgebrauch durchaus adäquat erscheinen. Der Stil der Dolmetschungen ist außerdem passend und auch die verwendete Lexik ist größtenteils entsprechend gewählt, sodass durchaus hervorzuheben ist, dass die nötige Loslösung von Formulierungen des Ausgangstextes stattfand. Allerdings finden sich ebenso einige Formulierungen in den Dolmetschungen, die lexikalisch nicht korrekt sind, wobei es sich auch hier stets um Fälle handelt, die keine inhaltlichen Widersprüche im Vergleich zum Ausgangstext produzieren.

Insgesamt handelt es sich demnach bei diesen Leistungen der Dolmetscher:innen um weitestgehend adäquate Dolmetschungen, wobei bei jeder Einzelnen individuelle Abstriche gemacht werden müssen. Im Gesamten betrachtet ergibt sich allerdings ein gutes Bild davon, welche Qualität eine menschliche Dolmetschung in diesem Kontext haben kann. Daraus ergeben sich nun die Möglichkeiten und Grenzen von Google Übersetzer in Bezug auf eine maschinelle Konsektivdolmetschung, worauf in weiterer Folge genauer eingegangen werden soll, um die eingangs erläuterte Forschungsfrage danach beantworten zu können.

6 Diskussion und Schlussfolgerungen

Ausgehend von den in den vorangegangenen Analysen der einzelnen Schritte der maschinellen Dolmetschung, sowie der insgesamt Bewertung dieser und jenen Leistungen der Humandolmetscher:innen soll im Rahmen dieses Kapitels nun die Diskussion der Ergebnisse mit dem Ziel, die Möglichkeiten und Grenzen der maschinellen Dolmetschung erläutern und damit die eingangs erwähnte Forschungsfrage beantworten zu können, erfolgen.

Zusammenfassend lässt sich nun feststellen, dass Qualität in der Dolmetschwissenschaft ein recht komplexes Konzept ist, zu dem es keine eindeutige oder allgemeingültige Definition gibt. Das liegt unter anderem daran, dass unterschiedlichste Erwartungen an Dolmetschleistungen erfüllt werden müssen, wodurch sich eine Vielzahl an Qualitätskriterien ergibt, die in verschiedenen Kontexten unterschiedlich gewichtet werden können. Dementsprechend wurde für die Analyse im Rahmen dieser Masterarbeit ausgehend von zuvor erarbeiteten Qualitätskriterien und möglichen Fehlertypen ein Bewertungsschema erstellt, dass eine Dolmetschbewertung grob in die folgenden zwei Bereiche gliedert: die Übertragung des Inhaltes und die Realisierung der Dolmetschung. Demzufolge wurden die Treue zum Ausgangstext, die Kohärenz, die Pragmatik, die Terminologie, die Stimme, das Sprechtempo, die Rhythmik, die Aussprache, die Grammatik, die Syntax, die Lexik, der Stil, die Dauer, die Intonation und die Prosodie als Parameter zur Bewertung herangezogen.

Da eine maschinelle Dolmetschung aus den drei aufeinanderfolgenden Schritten der Spracherkennung, der maschinellen Übersetzung und der Sprachausgabe besteht, wurden diese einzelnen Schritte auch einzeln beschrieben und analysiert, um genau herausfinden zu können, welche Herausforderungen und Problematiken sich an welcher Stelle einer maschinellen Dolmetschung ergeben. Dafür wurde der Ausgangstext dem maschinellen Übersetzungsprogramm Google Übersetzer vorgespielt, das daraufhin ein Transkript anfertigt, welches anschließend übersetzt und vorgelesen werden kann. Dadurch konnten das originalsprachliche Transkript, die Übersetzung und die Sprachausgabe jeweils einzeln anhand des eben beschriebenen Bewertungsschemas analysiert werden. Als Vergleichsbasis dienten die Dolmetschleistungen von vier Humandolmetscher:innen, die den Ausgangstext ebenso konsekutiv dolmetschten, wobei die Ergebnisse zunächst als Audiodateien vorlagen, die anschließend von mir transkribiert wurden.

Bei der Bewertung der maschinellen Dolmetschung sticht hervor, dass sich vor allem hinsichtlich der Realisierung der Dolmetschung einige Schwachstellen aufgrund der künstlichen, monotonen Stimme und der fehlenden Kürzungen ergeben. Hinzu kommen die fehlenden

pragmatischen Anpassungen, sowie einige doch recht gravierende Sinnverfälschungen aufgrund von Fehlern im Zuge der Spracherkennung, sodass der Zieltext teilweise inhaltliche Gegensätze zum Ausgangstext und gänzlich unverständliche Passagen enthält.

Bei den Leistungen der Humandolmetscher:innen hingegen ergaben sich Herausforderungen und Schwierigkeiten anderer Art, während insgesamt keine wirklich groben Sinnverfälschungen und damit gegensätzliche Aussagen aufzufinden sind. So sind vor allem einige Auslassungen und leichte grammatikalische und syntaktische Fehler zu erkennen, wobei andererseits – neben der korrekten Übertragung des Inhaltes – auch die Realisierung der Dolmetschung in Bezug auf Stimme, Prosodie und Intonation positiv hervorzuheben ist und einen großen Qualitätsunterschied zu der maschinellen Dolmetschung zeigt.

Genauer gesagt, fällt beim Vergleich zwischen Mensch und Maschine nun auf, dass den für die Rezipient:innen wohl offensichtlichster Unterschied die Stimme, mit der der Zieltext artikuliert wird, darstellt. So kann auf den ersten Blick vermutlich nicht direkt zwischen einer schriftsprachlichen Übersetzung eines Übersetzungsprogramms und eines Menschen unterschieden werden, vor allem wenn es sich um nahezu fehlerlose Passagen handelt – wie es abschnittsweise auch hier der Fall war – und der Ausgangstext nicht bekannt beziehungsweise verständlich ist – was bei Rezipient:innen einer Übersetzung oder Dolmetschung den Normalfall darstellt, sodass nicht von grammatikalischen Fehlern oder inhaltlichen Unterschieden auf eine maschinelle Übersetzung geschlossen werden kann. Bei einer Dolmetschung, bei welcher der Text allerdings in mündlicher Form präsentiert wird, steht die Stimme im Vordergrund und stellt einen wichtigen Punkt hinsichtlich der Erwartungen der Nutzer:innen und der Bewertungen von Dolmetschungen dar, wie in Kapitel 2.2 und 2.3 genauer erläutert wurde. Bei der Sprachausgabe von Google Übersetzer handelt es sich allerdings um eine recht künstlich wirkende Stimme, die von möglichen Rezipient:innen wohl recht schnell als nicht-menschliche Stimme erkannt werden sollte. Dies wiederum kann eine negative Bewertung einer solchen Dolmetschung nach sich ziehen, wenn bedacht wird, dass eine lebendige Stimme mit ausgeprägter Intonation erwartet wird und künstliche, monotone Stimmen als möglicher Störfaktor hinsichtlich einer qualitativ hochwertigen Dolmetschung gelten (vgl. Collados Aís 2016: 684f.).

In diesem Punkt lässt sich somit eine klare Grenze der Möglichkeiten von Google Übersetzer feststellen, da die Präsentation des gedolmetschten Zieltextes nicht an jene der Humandolmetscher:innen heranreicht und hier demnach ein eindeutiger Unterschied zu erkennen ist. Dabei ist nicht auszuschließen, dass sich die Technologie in diesem Punkt in Zukunft weiter entwickeln wird, wodurch es vielleicht künftig die Möglichkeit geben wird, dass sich auch synthetisch erzeugte Stimmen echten menschlichen Stimmen insofern annähern (vgl. Härtel 2016:

98f.), dass kein so eindeutiger Unterschied mehr festgestellt werden kann, wodurch diese aktuell noch sehr klare Grenze künftig vielleicht etwas verschwimmen und die Qualität maschineller Dolmetschungen in diesem Punkt an jene von Menschen herankommen könnte.

Außerdem schränkt die Funktionsweise der Spracherkennung die Qualität der anschließenden maschinellen Übersetzung stark ein. So lassen sich die meisten Fehler, die eine starke inhaltliche Verzerrung, eine Veränderung der Kohärenz oder die Unverständlichkeit einzelner Passagen des Zieltextes verursachen und damit die angestrebte Wirkung des Textes verfehlen, auf Fehler bei der Erstellung des Transkriptes durch die Spracherkennung von Google Übersetzer zurückführen, wie die Analyse zeigte. Dies ist dabei sicherlich stark abhängig von der jeweiligen Ausgangssprache. In diesem Fall stellt Französisch eine Sprache dar, die viele Homophone enthält, die für Menschen durch den Kontext heraus leicht ihrer jeweiligen Bedeutung zugeordnet werden können. Eine Spracherkennungssoftware hingegen kann in vielen Fällen den textuellen Kontext nicht erkennen, wodurch es passieren kann, dass ein falsches Wort erkannt und verschriftlicht wird, was in weiterer Folge zu einer Fehlübersetzung führen kann, wie auch die vorliegende Analyse gezeigt hat. Abgesehen von solchen sprachspezifischen Homophonen können auch Eigenheiten der jeweiligen Sprecher:innen oder Umgebungsgeräusche die Funktionsweise der Spracherkennungssoftware einschränken. So können Umgebungsgeräusche die Sprecher:innen übertönen oder Füllwörter, Verzögerungslaute, Husten, Räuspern oder Ähnliches seitens der Sprecher:innen können von der Spracherkennungssoftware als vermeintlich bedeutungstragende Lexeme aufgefasst und verschriftlicht werden, was wiederum zu Fehlern in der maschinellen Dolmetschung führt.

Abgesehen von jenen inhaltlichen Verzerrungen aufgrund der fehlerhaften Spracherkennung kam es im Zuge des Schrittes der maschinellen Übersetzung zu weiteren Problematiken, die zusätzliche Grenzen des Prozesses der maschinellen Dolmetschung aufzeigen. So finden sich einige für das Deutsche unpassende oder ungewöhnliche Formulierungen, die etwa unterschiedliche Konnotationen in Ausgangs- und Zielsprache aufweisen oder auch zu wörtliche Übersetzungen. Diese stellen aber – im Gegensatz zu jenen Problemen, die durch die Spracherkennung verursacht wurden – keine groben Sinnveränderungen dar, sodass die Rezipient:innen sich daran zwar stören könnten, aber den Inhalt trotzdem verstehen würden. Hinzu kommt, dass es im Vergleich zum französischen Transkript von Google Übersetzer sogar eine Verbesserung bezüglich Satzzeichensetzung gibt, sodass im Zuge der maschinellen Übersetzung teilweise kleinere Fehler der Spracherkennungssoftware ausgebessert werden konnten.

Im Schritt der Sprachausgabe hingegen sind keine groben Fehler passiert, wobei auch hier anzumerken ist, dass der Text keine großen Herausforderungen, wie etwa Homographie mit

unterschiedlicher Aussprache, Eigennamen oder Fremdwörter enthält. Dementsprechend konnte durch die Sprachausgabe zwar der Inhalt gut übermittelt werden, allerdings müssen deutliche Abstriche bezüglich der bereits erwähnten Qualität des Klanges der Stimme, der Prosodie und der Intonation gemacht werden. Hinsichtlich der Prosodie und Intonation ist festzuhalten, dass die Sprachausgabe in einem ähnlichen Tempo wie jenes der Humandolmetscher:innen erfolgt, was positiv hervorzuheben ist. Dazu kommt, dass die Aussprache klar und deutlich und in einem guten Rhythmus erfolgt. Ein weiterer Vorteil einer synthetischen Stimme ist, dass diese keine Verzögerungslaute produziert, die bei den menschlichen Dolmetschleistungen teilweise vorhanden waren. Auch werden keine zusätzlichen Fülllaute produziert, sodass nur jene vorkommen können, die die Spracherkennungssoftware aus dem Ausgangstext erkannt und verschriftlicht hat. Allerdings erfolgen die Sprechpausen an teilweise unnatürlichen Stellen, während diese an anderen Stellen wiederum fehlen, was als Störfaktor seitens der Nutzer:innen angesehen werden kann. Außerdem könnten solche deplatzierten Pausen nicht nur unnatürlich klingen, sondern im schlimmsten Fall sogar zu Verständnisproblemen führen, wenn diese etwa mitten in einer Phrase erfolgen oder wenn Pausen zwischen unterschiedlichen Sinneinheiten fehlen. Zusätzlich klingt die synthetische Stimme von Google Übersetzer über weite Strecken recht monoton, da wenig Intonation stattfindet, was ebenso als Störfaktor zu notieren ist, da monotones Sprechen zu einer schnelleren Ermüdung und damit einhergehendem schlechterem inhaltlichem Verständnis bei den Rezipient:innen führen kann (vgl. Collados Aís 2016: 685).

Insgesamt lässt sich nun zu den Möglichkeiten und Grenzen der Funktion zur maschinellen Dolmetschung der App Google Übersetzer festhalten, dass diese im Vergleich zu Humandolmetscher:innen insgesamt deutlich mehr Fehler in sich birgt und diese Fehler auch als gravierender zu erachten sind als jene Fehler, die bei den Humandolmetscher:innen auftauchen. So finden sich im Zieltext der maschinellen Dolmetschung beispielsweise grobe inhaltliche Widersprüche im Vergleich zum Ausgangstext, wodurch eine starke Sinnverfälschung stattfindet und das eigentliche Ziel der Dolmetschung nicht erreicht werden kann. Im Gegensatz dazu handelt es sich bei den Fehlern der Humandolmetscher:innen nahezu ausschließlich um Fehler, die für die Rezipient:innen teilweise zwar wohl leichter erkennbar sind, wenn es sich etwa um grammatikalische Fehler handelt, oder die inhaltliche Auslassungen und damit Informationsverlust mit sich bringen, aber insgesamt dennoch keine so starken Auswirkungen auf die korrekte Übertragung des Sinnes des Ausgangstextes für das Zielpublikum haben.

So finden sich bei den menschlichen Dolmetschleistungen zwar Kürzungen von Redundanzen und auch einige weitere Stellen, an denen zusammengefasst wurde oder einige inhaltliche Details weggelassen wurden. Jedoch findet in allen vier Fällen eine korrekte Übertragung

des Inhaltes statt. Selbiges gilt auch für die vorkommenden Sinnverfälschungen. Auch diese beziehen sich folglich fast ausschließlich auf stilistische Nuancen oder inhaltliche Details, so dass Abänderungen an diesen Stellen zu keinen großen inhaltlichen Einbußen führen.

Diese Erkenntnis aus der qualitativen Analyse lässt sich auch mit den Ergebnissen der quantitativen Analyse bestätigen. Auch hier stechen – wie auch in der nachfolgenden Tabelle verdeutlicht wird – vor allem die zahlreichen als schwere Fehler eingestuften Mängel an der Dolmetschleistung von Google Übersetzer hervor, während solche bei den Humandolmetscher:innen nahezu gar nicht vorkommen. Ähnliches gilt auch für mittlere Fehler, sodass auch diese häufiger bei der maschinellen Dolmetschung vorkommen. Lediglich bei den leichten Fehlern finden sich bei den menschlichen Dolmetschleistungen mehr, die vor allem auf die oben angesprochenen Auslassungen und leichten Grammatik- und Syntaxfehler zurückzuführen sind. Laut der quantitativen Fehleranalyse finden sich bei der Dolmetschung von Google Übersetzer insgesamt ein leichter, sechs mittlere und neun schwere Fehler in Bezug auf die Übertragung des Inhalts und vier leichte, fünf mittlere und ein schwerer Fehler in Bezug auf die Realisierung der Dolmetschung. Bei den Humandolmetschungen hingegen ergeben sich Durchschnittswerte der vier Dolmetschleistungen von 4,75 leichten Fehlern, 3,5 mittleren Fehlern und 0,5 schweren Fehlern bezüglich der Übertragung des Inhalts und fünf leichten und 3,75 mittleren Fehlern bezüglich der Realisierung der Dolmetschung. Insgesamt ergeben sich dadurch fünf leichte Fehler, elf mittlere Fehler und zehn schwere Fehler bei der Dolmetschung von Google Übersetzer und 9,75 leichte Fehler, 7,25 mittlere Fehler und 0,5 schwere Fehler bei dem Durchschnitt der menschlichen Dolmetschleistungen.

Dementsprechend lässt sich feststellen, dass von den Humandolmetscher:innen zwar mehr leichte Fehler in Bezug auf die Übertragung des Inhalts gemacht wurden, diese jedoch das Verständnis der Rezipient:innen nur leicht stören könnten. Im Gegensatz dazu wurden von Google Übersetzer nicht nur mehr mittlere Fehler in Bezug auf die Übertragung des Inhalts gemacht, sondern vor allem deutlich mehr schwere Fehler begangen. Diese beziehen sich dabei sowohl auf die Übertragung des Inhalts als auch auf die Realisierung der Dolmetschung. Solche schweren Fehler kommen bei dem Durchschnitt der Leistungen der menschlichen Dolmetscher:innen nahezu gar nicht vor. Diese schweren Fehler sind insofern relevant, als dass diese das Verständnis des Inhaltes deutlich erschweren oder sogar unmöglich machen und damit die Qualität der Dolmetschung deutlich verschlechtern, da die Verständlichkeit des Inhaltes für die Rezipient:innen die grundlegende Basis für eine qualitativ hochwertige Dolmetschleistung darstellt.

Zur Verdeutlichung und für eine bessere Übersicht ist diese Gegenüberstellung der Anzahl der Fehler von Google Übersetzer beziehungsweise dem Durchschnitt der Anzahl der Fehler der menschlichen Dolmetscher:innen auch in der folgenden Tabelle aufgeführt:

Hauptkategorie Google Übersetzer	Fehlerkategorie	Leicht	Mittel	Schwer	Gesamt
Übertragung des Inhalts	Treue zum Ausgangstext	1	4	7	12
	Kohärenz	0	0	2	2
	Pragmatik	0	2	0	2
Realisierung der Dolmetschung	Grammatik	1	3	1	5
	Lexik	3	2	0	5
Gesamt Google Übersetzer		5	11	10	26
Hauptkategorie Dolmetscher:innen	Fehlerkategorie	Leicht	Mittel	Schwer	Gesamt
Übertragung des Inhalts	Treue zum Ausgangstext	4,75	2,75	0,5	8
	Pragmatik	0	0,75	0	0,75
Realisierung der Dolmetschung	Grammatik	1,5	1,5	0	3
	Syntax	2	0,25	0	2,25
	Lexik	1,5	2	0	3,5
Gesamt Dolmetscher:innen		9,75	7,25	0,5	17,5

Vergleich der Fehleranalyse zwischen Google Übersetzer und den Dolmetscher:innen

Insgesamt ergibt sich nun, dass sich bei einer maschinellen Dolmetschung mit der App Google Übersetzer bei allen drei Schritten Probleme ergeben können, sodass die Erwartungen der Rezipient:innen nicht vollends eingehalten werden können. Andererseits funktioniert die Sinnübertragung über weite Strecken relativ problemfrei, sodass der Einsatz einer solchen Software wohl je nach Einsatzgebiet für den konkreten Fall abgewogen werden kann. Für den privaten Gebrauch für einfache, kurze Alltagsgespräche kann eine solche Software sicherlich recht gut eingesetzt werden. Geht es jedoch um inhaltlich und terminologisch komplexere und sensiblere Themen in einem professionellen Kontext, in dem keine inhaltliche Fehler passieren dürfen, stößt der Einsatz von Google Übersetzer wohl recht schnell an eine Grenze, wo ein deutlicher Unterschied zu Humandolmetscher:innen auszumachen ist, der vor allem durch das Kontextwissen und das Verständnis des wiedergegebenen Inhalts bedingt ist.

Nichtsdestotrotz ist festzuhalten, dass eine solche maschinelle Dolmetschleistung in Bezug auf den finanziellen und organisatorischen Aufwand deutlich weniger Ressourcen benötigt, als dies bei einer professionellen menschlichen Dolmetschleistung der Fall ist. Dementsprechend würde die hier untersuchte maschinelle Dolmetschung etwa durchaus für eine Situation ausreichen, in der etwa finanzielle Mittel fehlen oder es sich um einen sehr spontan benötigten Dolmetscheinsatz handelt und gleichzeitig die Anforderungen an die Zieltextproduktion nicht allzu hoch sind.

So kann der hier untersuchte Zieltext etwa sehr wohl dafür verwendet werden, damit Rezipient:innen, die der Ausgangssprache nicht mächtig sind, einen ersten Eindruck des Inhalts des Ausgangstextes bekommen, den diese ansonsten gar nicht verstehen würden. So könnte das hier untersuchte Programm etwas für Alltagskommunikationssituationen, die über Sprachgrenzen hinaus geschehen, eingesetzt werden, um eine grundlegende Kommunikation zu ermöglichen und einen Überblick über die Inhalte zu gewähren, während gerade Details mit Vorsicht zu genießen sind.

Allerdings sind aufgrund der oben erläuterten Fehler und Problematiken, die sich bei dieser maschinellen Dolmetschung ergaben, auch deutliche Grenzen in Bezug auf den Einsatzbereich einer solchen Software zu setzen. So kann etwas bei einem Einsatz in einer professionellen Kommunikationssituation, in der komplexe Themen und sensible Inhalte besprochen werden, wo inhaltliche Fehler gravierende Konsequenzen mit sich bringen würden, nicht garantiert werden, dass alle Inhalte des Ausgangstextes auch korrekt in den Zieltext übertragen werden. Programme für maschinelle Übersetzung können beispielsweise keine situativen Anpassungen vornehmen, da ihnen das Kontextwissen fehlt.

Außerdem stellt sich gerade bei sensiblen Themen und Daten auch die Frage des Datenschutzes, da bei der App Google Übersetzer nicht direkt ersichtlich ist, was mit den eingespielten Daten passiert und wo diese gespeichert und in weiterer Folge verwendet werden. Dies kann bei menschlichen Dolmetscher:innen etwa mit einer Verschwiegenheitserklärung gelöst werden, in der festgehalten wird, dass keine inhaltlichen Informationen über den zu dolmetschenden Text mit jeglichen Außenstehenden geteilt werden dürfen.

Dabei ist allerdings ebenso zu beachten, dass die vorliegende Analyse lediglich einen Einblick in die aktuelle Funktionsweise von Google Übersetzer bietet, während es weitere Programme auf dem aktuellen Markt gibt, die eventuell in einigen Punkten andere Ergebnisse erzielen würden. Außerdem wird stets weiter an Systemen zur Spracherkennung, maschinellen Übersetzung und Sprachausgabe geforscht, weswegen in den nächsten Jahren wohl durchaus neue und verbesserte Programme veröffentlicht werden können, weswegen es gilt, sich auch zukünftig mit diesem hochaktuellen Forschungsfeld zu befassen, um Weiterentwicklungen und potentielle Verbesserungen der Programme für maschinelles Dolmetschen im Blick behalten zu können.

Demnach bleibt abzuwarten, wohin eine Verbesserung beziehungsweise eine Weiterentwicklung der verwendeten Technologien führen kann. So wäre es beispielsweise wünschenswert, die Realisierung einer solchen maschinellen Dolmetschung zu verbessern, sodass etwa der Klang der synthetischen Stimme derart verbessert wird, dass dieser mehr einer menschlichen Stimme gleicht und es demnach für Rezipient:innen angenehmer wird, der Dolmetschung zuzuhören. Dazu sollte auch gleichzeitig die Intonation der synthetischen Stimme, sowie deren Pausenverhalten angepasst werden, da auch bei diesen Punkten noch deutliche Unterschiede zwischen Mensch und Maschine auszumachen sind.

Zusätzlich müsste aber klarerweise auch – wie oben bereits angesprochen wurde – dahingehend weitergearbeitet werden, dass auch die beiden anderen Schritte – die Spracherkennung und die maschinelle Übersetzung – stetig weiterentwickelt werden, da sich hier die größten Fehlerquellen hinsichtlich einer adäquaten Übertragung des Inhaltes des Ausgangstextes in die Zielkultur und damit einer qualitativ hochwertigen Dolmetschung ergeben, wie auch die Analyse zeigte.

Abgesehen von der Untersuchung zukünftiger Versionen oder weiteren Programmen zur maschinellen Dolmetschung stellten sich im Verlauf dieser Arbeit weitere Fragen, die in zukünftigen Forschungsarbeiten behandelt und beantwortet werden könnten. Demnach sei abermals angemerkt, dass es sich bei dem Text, der in der vorliegenden Arbeit einer maschinellen Dolmetschung unterzogen wurde, um eine idealisierte Version handelt, die so in der

Praxis wohl kaum vorkommt, da keine Störgeräusche vorkommen, die Sprecherin recht klar und deutlich spricht, es zu keinen technischen Problemen, wie etwa Mikrofonstörungen oder ähnlichem kommt, und auch der Inhalt und die verwendete Terminologie recht einfach gehalten sind. Dadurch konnten einige mögliche Fehlerquellen schon im Vorfeld ausgeschaltet werden, was hier bewusst gemacht wurde, um sich rein auf die Dolmetschung an sich zu konzentrieren, ohne zusätzliche Problematiken beachten zu müssen. Wird nun allerdings eine solche maschinelle Dolmetschung in einer realen Situation angewandt, müssen diese Herausforderungen ebenso bedacht werden, weshalb die Dolmetschleistung in einer solchen Situation qualitativ durchaus schlechter ausfallen könnte. So wäre es sicherlich interessant, weitere Herausforderungen für eine solche Software miteinzubeziehen, die aufbauend auf den Ergebnissen dieser Arbeit weitere und komplexere Thematiken behandeln. Dazu könnten etwa in Hinblick auf die Spracherkennung zusätzliche Störgeräusche, Eigennamen, Fremdwörter, Neologismen, Füllwörter, Verzögerungslaute und Selbstkorrekturen der Sprecher:innen oder auch Sprecher:innenwechsel zählen. Auch für den Schritt der maschinellen Übersetzung könnten mit der Anforderung von stilistischen und pragmatischen Anpassungen, sowie schwieriger zu übersetzenden Textpassagen, die textuelles, kulturelles oder situatives Kontextwissen erfordern, zusätzliche Schwierigkeiten geschaffen werden, deren Umsetzung durch ein Programm zur maschinellen Dolmetschung sicherlich aufschlussreich zu erforschen wäre. Für den Prozess der Sprachausgabe gilt Ähnliches wie für die Spracherkennung, sodass das Vorhandensein von Fremdwörtern, Abkürzungen, Neologismen, Eigennamen und syntaktischer oder lexikalischer Ambiguität als weiterer möglicher Schwachpunkt gelten kann und dementsprechend einen relevanten Untersuchungsgegenstand darstellen könnte.

Außerdem könnte es durchaus interessant sein, sich die Anhäufung von Fehlern gegen Ende des hier untersuchten Textes nochmals genauer anzusehen, um eventuell eruieren zu können, ob eine solche Software für maschinelles Dolmetschen grundsätzlich ein Problem mit länger dauernden Texten hat, oder ob es sich hier nur um einen Zufall handelt. Nicht zuletzt sei anzumerken, dass in dieser Arbeit lediglich der Modus des Konsekutivdolmetschens untersucht wurde, sodass für mögliche zukünftige Arbeiten jene erläuterten möglichen Forschungsziele auch mit einer Software durchgeführt werden könnte, die maschinelles Simultandolmetschen ermöglicht. In einem solchen Fall würden sich weitere spannende Fragen ergeben, deren Erforschung künftig sicherlich lohnenswert wäre.

Nichtsdestotrotz ist das Forschungsfeld der maschinellen Dolmetschung von großer Relevanz, da durch die stetigen Verbesserungen an den eingesetzten Technologien auch die Qualität der erzeugten Produkte immer weiter verbessert werden kann, weswegen es auch in

Zukunft gilt, weitere Forschungen mit verbesserten, weiterentwickelten oder neuen Programmen zu betreiben, um herausfinden zu können, welche Möglichkeiten und Grenzen sich in Zukunft bei dem Einsatz von Tools zur maschinellen Dolmetschung noch ergeben werden. Ein Beispiel wäre hier etwa die Weiterentwicklung von Programmen für eine maschinelle Simultandolmetschung, da über Google Übersetzer in der Version, die in dieser Arbeit verwendet wurde, etwa nur eine maschinelle Konsektivdolmetschung möglich ist, deren Einsatzmöglichkeiten im Bereich des Konferenzdolmetschens im Vergleich zu Simultandolmetschungen aufgrund etwa der längeren Dauer etwas eingeschränkt sind. Es bleibt also abzuwarten, wohin sich die Technologien in der Zukunft noch entwickeln werden und welche neuen Möglichkeiten der maschinellen Translation dadurch entstehen können.

Bibliographie

- Ahrens, Barbara (2016). Konsektivdolmetschen. In: Kadrić, Mira & Kaindl, Klaus (Hg). *Berufsziel Übersetzen und Dolmetschen; Grundlagen, Ausbildung, Arbeitsfelder*. Tübingen: A. Francke, 84-102.
- Bühler, Hildegund (1986). Linguistic (semantic) and extra-linguistic (pragmatic) criteria for the evaluation of conference interpretation and interpreters. *Multilingua* 5 (4), 231-235.
- Champoux, Ménaïc (2018). L'orthographe des homophones : quelles difficultés pour les élèves du secondaire ?. *Bellaterra Journal of Teaching & Learning Language & Literature* 11 (1), 63-64.
- Christensen, Tina Paulsen (2011). User expectations and evaluation: a case study of a court interpreting event. *Perspectives* 19 (1), 1-24.
- Colina, Sonia (2011). Evaluation/Assessment. In: Gambier, Yves & van Doorslaer, Luc (Hg.) *Handbook of Translation Studies; Volume 2*. Amsterdam/Philadelphia: John Benjamins, 43-48.
- Collados Aís, Ángela (2016). La entonación monótona y la calidad de la interpretación simultánea: frecuencia, conceptualizaciones y efectos. *Meta* 61 (3), 675-691.
- Cürten, Giulia (2016). *Maschinelles Dolmetschen mit Google Übersetzer; Eine Analyse im Sprachenpaar Deutsch / Italienisch*. Masterarbeit, Universität Wien.
- Dam, Helle V. (2010). Consecutive interpreting. In: Gambier & van Doorslaer (Hg.), 75-79.
- Deng, Li & O'Shaughnessy, Douglas (2003). *Speech processing. A dynamic and optimization-oriented approach*. London/New York: Taylor & Francis Group.
- Fantinuoli, Claudio (2025). Machine interpreting. In: Davitti, Elena & Korybski, Tomasz & Braun, Sabine (Hg.). *The Routledge handbook of interpreting, technology and AI*. London/New York: Routledge.
- Fantinuoli, Claudio & Prandi, Bianca (2021). Towards the evaluation of automatic simultaneous speech translation from a communicative perspective. *Proceedings of the 18th international conference on spoken language translation*, 245-254.
- Forcada, Mikel L. (2010). Machine translation today. In: Gambier & van Doorslaer (Hg.), 215-223.

- Fünfer, Sarah (2013). *Mensch oder Maschine: Dolmetscher und maschinelles Dolmetschsystem im Vergleich*. Berlin: Frank & Timme.
- Gambier, Yves & van Doorslaer, Luc (Hg.) (2010). *Handbook of translation studies; Volume 1*. Amsterdam / Philadelphia: John Benjamins.
- Gile, Daniel (2001). L'Évaluation de la qualité de l'interprétation en cours de formation. *Meta* 46 (2), 379-393.
- Gile, Daniel (2009). *Basic concepts and models for interpreter and translator training*. Amsterdam / Philadelphia: John Benjamins.
- Glocknitzer, Julia (2020). *Arzt-Patient-Kommunikation mit SayHi Translate. Eine aufgabenbezogene Evaluierung im Sprachenpaar Deutsch-Ungarisch*. Masterarbeit, Universität Wien.
- Gouadec, Daniel (2010). Quality in translation. In: Gambier & van Doorslaer (Hg.), 270-275.
- Grbić, Nadja (2008). Constructing interpreting quality. *Interpreting* 10 (2), 232-257.
- Han, Chao (2018). Using rating scales to assess interpretation: practices, problems and prospects. *Interpreting* 20 (1), 59-95.
- Hansen, Gyde (2010). Translation 'errors'. In: Gambier & van Doorslaer (Hg.), 385-388.
- Härtel, Johannes (2016). Vollautomatisiertes Dolmetschen. *Lebende Sprachen* 61 (1), 67-116.
- Kalina, Sylvia (2004). Zum Qualitätsbegriff beim Dolmetschen. *Lebende Sprachen* 49 (1), 2-8.
- Kalina, Sylvia (2012). Quality in interpreting. In: Gambier, Yves & van Doorslaer, Luc (Hg.). *Handbook of Translation Studies; Volume 3*. Amsterdam / Philadelphia: John Benjamins, 134-140.
- Koehn, Philipp (2020). *Neural machine translation*. Cambridge: Cambridge University Press.
- Kopczyński, Andrzej (1984). Problems of quality in conference interpreting. In: Fisiak, Jacek (Hg.). *Contrastive Linguistics*. Berlin & New York & Amsterdam: Mouton, 283-300.
- Kopczyński, Andrzej (1994). Quality in conference interpreting: Some pragmatic problems. In: Lambert, Sylvie & Moser-Mercer, Barbara (Hg.). *Bridging the Gap; Empirical Research in simultaneous interpretation*. Amsterdam/Philadelphia: John Benjamins, 87-100.

- Kurz, Ingrid (2001). Conference interpreting: quality in the ears of the users. *Meta* 46 (2), 394-409.
- Kutz, Wladimir (2005). Zur Bewertung der Dolmetschqualität in der Ausbildung von Konferenzdolmetschern. *Lebende Sprachen* 50 (1), 14-34.
- Leitner, Julia (2020). *Maschinelles Dolmetschen mit iTranslate; Ein Vergleich zwischen Mensch und Maschine*. Masterarbeit, Universität Wien.
- Moser, Peter (1995). *Survey on expectations of users of conference interpreting*. Genf: AIIC.
- Pfister, Beat & Kaufmann, Tobias (2008). *Sprachverarbeitung. Grundlagen und Methoden der Sprachsynthese und Spracherkennung*. Berlin/Heidelberg: Springer.
- Pöchhacker, Franz (2001). Quality assessment in conference and community interpreting. *Meta* 46 (2), 410-425.
- Pöchhacker, Franz (2004). *Introducing interpreting studies*. London/New York: Routledge.
- Pöchhacker, Franz (2010). Interpreting. In: Gambier & van Doorslaer (Hg.), 153-157.
- Pöchhacker, Franz (2015). User expectations. In: Pöchhacker, Franz (Hg.). *Routledge encyclopedia of interpreting studies*. London/New York: Routledge, 430-432.
- Pöchhacker, Franz (2024). Is machine interpreting interpreting?. *Translation Spaces*. <https://www.jbe-platform.com/content/journals/10.1075/ts.23028.poc> (Stand: 03.06.2025)
- Poibeau, Thierry (2017). *Machine translation*. Cambridge, MA: The MIT press.
- Rehbein, Jochen (2004). *Handbuch für das computergestützte Transkribieren nach HIAT: Version 1.0*. Hamburg: Sonderforschungsbereich 538.
- Ruske, Günther (1994²). *Automatische Spracherkennung. Methoden der Klassifikation und Merkmalsextraktion*. München/Wien: Oldenbourg.
- Setton, Robin (2010). Conference interpreting. In: Gambier & van Doorslaer (Hg.), 66-74.
- Shlesinger, Miriam/Déjean le Féal/Karla, Kurz, Ingrid/Cattaruzza, Lorella/Mack, Gabriele/Nilsson, Anna-Lena/Niska, Helge & Pöchhacker, Franz & Viezzi, Maurizio (1997). Quality in Simultaneous Interpreting. In: Gambier, Yves & Gile, Daniel & Taylor, Christopher (Hg.). *Conference interpreting; current trends in research*. Amsterdam/Philadelphia: John Benjamins, 123-131.

- Spilker, Jörg/Klarner, Martin & Görz, Günther (2000). Processing of Self-Corrections in a Speech-to-Speech System. In: Wahlster, Wolfgang (Hg.). *Verbmobil. foundations of speech-to-speech translation. Artificial intelligence*. Berlin/New York: Springer, 131-140.
- Taylor, Paul (2009). *Text-to-speech synthesis*. Cambridge: University Press.
- Turovsky, Barak (2016). Found in translation: more accurate, fluent sentences in Google Translate. <https://blog.google/products/translate/found-translation-more-accurate-fluent-sentences-google-translate/> (Stand: 19.03.2025).
- Zhang, Jingyi/Utiyama, Masao/Sumita, Elichro & Neubig, Graham & Nakamura, Satoshi (2020). Improving neural machine translation trough phrase-based soft forced decoding. *Machine Translation* 34, 21-39.

Anhang

T 1 – Transkript 1 – Selbst erstelltes Transkript des Ausgangstextes

1 Bonjour, je vous parle de la te/ technologie de la reconnaissance • faciALE et • surtout de son
2 utilisation répandue • en Chine. • • Je commence. • • Vous avez certainement déjà tous
3 entendu parler de:: la technologie de la reconnaissance faciale. • Il s'agit d'une technologie qui
4 permet ((ea)) de reconnaître une personne grâce à son visage ((ea)) et ceci de manière
5 complètement automatique. ((2s)) Et • • pourquoi est-ce qu'on la connaît bien? Tout
6 simplement parce que cette technologie commence • aujourd'hui à s'immiscer ((ea)) dans
7 notre quotidien. ((ea)) Alors, pour objectif • de faciliter la vie des citoyens.

8
9 • • Prenons l'exemple/ par exemple: euh de::s téléphones portables que l'on peut déverrouiller
10 ((ea)) grâce • à la reconnaissance faciale. ((ea)) Cela permet à l'utilisateur de ne pas devoir
11 reconnaî/ de n'pas devoir se souvenir d'un mot de passe compliqué ((ea)) et de ne pas • avoir
12 euh l'entrer/ l'encoder chaque fois qu'il veut • utiliser son téléphone. ((2s)) Et quand on parle
13 de nouvelles technologies e::t de leur application, ((ea)) on est • souvent forcé • de faire
14 référence à la Chine. • • ((ea)) En effet, la Chine • est souvent très friANDE euh des avancées
15 technologiques. • • ((ea)) Et l'exemple de la reconnaissance faciale ne fait PAS • exception. • •

16
17 ((ea)) La reconnaissance faciale est en effet déjà confortablement installée ((ea)) dans le
18 quotidien • des citoyens chinois. • ((ea)) Elle est utilisée par exemple euh aux caisses • des
19 supermarchés pour • payer ses courses. • • ((ea)) E:t les Chinois en généra/ en général
20 semblent assez favorables à:/au recours à cette technologie ((ea)) même si • elle peut paraître
21 parfois un peu intrusive. • • ((ea)) Selon eu::x, euh: la: reconnaissance faciale leur facilite euh
22 la vie e:t elle est euh: synonyme de nombreux bénéfices. ((ea)) En plus, euh selon eux, la
23 reconnaissance faciale permet • d'assurer la sécurité publique.

24
25 • • ((ea)) En effet, à l'origine, euh: la reconnaissance faciale en Chine était surtout utilisée
26 pour assurer la sécurité. ((ea)) • • Elle permet notamment de: /d'accélérer les contrôles
27 d'identité • dans les transports. Elle permet aussi • d'identifier des citoyens désobéissants. • •
28 ((ea)) Toutefois •, récemment •, une mesu:re prise par le gouvernement chinois semble • aller
29 trop loin. • ((ea)) Elle a suscité l'inquiétude voire la révolte de certains citoyens.

30
31 • • ((ea)) Le gouvernement a: décidé de rendre la reconnaissance faciale OBLigatoire dans
32 certains cas. ((ea)) Il s'agit plus précisément • du: cas des personnes souhaitant acheter un
33 nouveau téléphone portable. ((ea)) Aujourd'hui, tous les Chinois qui veulent acheter un
34 nouveau téléphone portable • doivent • faire enregistrer leurs visages ((ea)) pour pouvoir être
35 reconnus grâce à: la reconnaissance faciale. ((ea)) La Chine veut en:/ en effet pourvoir euh:
36 identifier euh tous les utilisateurs grâce à leurs numéros de téléphone et à leurs visages. ((ea))
37 Alors, il faut savoir que cette obligation de: /d'enregistrer les citoyens sous leur IDentité
38 réelle, ((ea)) elle est déjà en vigueur en Chine depuis 2013.

39
40 • • ((ea)) Elle vise surtout euh: cette/ cette initiative à surveiller les réseaux sociaux. ((ea)) En
41 effet, la majorité des réseaux sociaux en Chine ((ea)) euh: ont l'obligatION de connaître
42 l'identité réELLE des utilisateurs. ((ea)) Mais ce qui est nouveau ici, c'est justement le recours
43 euh à euh la reconnaissance faciale. • • ((ea)) Alors, pourquoi est-ce que les citoyens chinois

44 sont particulièrement inquiets de cette mesure? ((ea)) Bien, surtout parce qu'ils • craignent que
 45 leurs données euh: ainsi enregistrées ((ea)) puissent être transmises et revendues à d'autres
 46 entreprises. • • Il effet/ Il est en effet pas rare en Chine ((ea)) de constater des fuites • de
 47 données personnelles • • des personnes ((ea)) malintentionnées qui • reVENDent • les données
 48 personnelles des utilisateurs ((ea)) à d'autres entreprises.
 49
 50 • • Alors, en général, ils sou/s'agit de données telles que les adresses ou les numéros de
 51 téléphone, ((ea)) mais dans ce cas-CI et ce qui inquiète plus particulièrement les citoyens,
 52 ((ea)) c'est que ça concerne leurs données euh biométriques. C'est-à-dire, ((ea)) leurs
 53 empreintes digitales, • les scans de leurs visages ou encore de l'iris de leur œil. • • ((ea)) Un
 54 rapport publié récemment a en effet indiqué qu'environ 10% des applications sur téléphones
 55 portables ((ea)) étaient soupçonnées de collecter • de manière excessive les données
 56 biométriques • de leurs utilisateurs.
 57
 58 ((2s)) Et récemment encore c'est un autre exemple qui:: semble témoigner de: la méfiance
 59 croissante des citoyens chinois. • • Un professeur qui voulait se rendre da:ns une réserve
 60 d'animaux non loin de Shanghai ((ea)) a euh: décidé euh de refuser • d'entrer dans la réserve.
 61 Pourquoi ? ((ea)) Parce que le:s gardiens à l'entrée voulaient scanner son visage. ((ea)) • • Il l'a
 62 refusé que son visage soit scanné • et il a ensuite décidé d'attaquer en justice • la direction
 63 d'une réserve d'animaux. ((ea)) Euh ce genre de:: de cas est tout à fait nouveau ((ea)) e:t euh
 64 sans précédent en Chine et montre donc que les citoyens en ont marre. • • ((ea)) Mais qu'est-
 65 ce qui pousse le gouvernement chinois à: imposer la reconnaissance faciale • dans • tous les
 66 aspects du quotidien? ((ea)) Alors officiellement, évidemment, comme je l'ai dit ((ea)) il s'agit
 67 d'une mesure visant à assurer la sécurité des ci/ des citoyens. ((ea)) Mais il faut savoir
 68 également que • le gouvernement chinois est désiREUX de faire de la Chine un leaDER
 69 technologique. ((ea)) • Par conséquent, ((ea)) • • le gouvernement n'hésite pas à offrir son
 70 soutien à de:s progRAMMES de reconnaissance faciale ou d'intelligence artificielle en
 71 général.
 72
 73 ((ea)) • • En Chine •, on compte enviro:n 850 millions d'utilisateurs • de téléphones
 74 portables. Ça constitue une euh: précieuse base de données ((ea)) pour les entreprises et les
 75 administrations. • ((ea)) Alors, pour tenter de rassurer la population ((ea)) le gouvernement a
 76 publié des directives • non contraignantes visant à réglementer la collecte et l'utilisation des
 77 données. ((ea)) Mais • alors actuelle la Chine ne dispose pas de VRAIES législations à ce
 78 sujet. • • Alors, • apparemment, une loi est en cours de rédaction. Mais • • suffira-t-elle à
 79 apaiser les chito/citoyens chinois ? ((ea)) • Bien, l'avenir nous dira. • Merci.

T 2 – Transkript 2 – Transkript des Ausgangstextes von Google Übersetzer

1 bonjour, je vous parle de la technologie de la reconnaissance faciale et surtout de son
2 utilisation répandue en Chine et je commence. vous avez certainement déjà tous entendu
3 parler de la technologie de la reconnaissance faciale, il s'agit d'une technologie qui permet de
4 reconnaître une personne grâce à son visage et ceci de manière complètement automatique. et
5 pourquoi est-ce qu'on la connaît bien tout simplement parce que cette technologie commence
6 aujourd'hui à s'immiscer dans notre quotidien elle a pour objectif de faciliter la vie des
7 citoyens.

9 Prenons l'exemple par exemple des téléphones portables que l'on peut déverrouiller grâce à la
10 reconnaissance faciale cela permet à l'utilisateur de ne pas devoir reconnu de ne pas devoir se
11 souvenir d'un mot de passe compliqué et de ne pas avoir à l'entrée l'encoder chaque fois qu'il
12 veut utiliser son téléphone. Et quand on parle de nouvelles technologies et de leur application,
13 on est souvent forcé de faire référence à la Chine en effet, la Chine est souvent très friande des
14 avancées technologiques et l'exemple de la reconnaissance faciale ne fait pas exception.

16 La reconnaissance faciale est en effet déjà confortablement installé dans le quotidien des
17 citoyens chinois. Elle est utilisée par exemple aux caisses des supermarchés pour payer ses
18 courses. et les Chinois en général en général semblent assez favorable à recours à cette
19 technologie même si elle peut paraître parfois un peu intrusive. selon eux la reconnaissance
20 faciale leur facilité la vie et elle est synonyme de nombreux bénéfices en plus selon eux la
21 reconnaissance faciale permet d'assurer la sécurité publique.

23 en effet à l'origine la reconnaissance faciale en Chine était surtout utilisée pour assurer la
24 sécurité. Elle permet notamment de d'accélérer les contrôles d'identité dans les transports, elle
25 permet aussi d'identifier des citoyens des obéissants. toutefois récemment une mesure prise
26 par le gouvernement chinois semble aller trop loin, elle a suscité l'inquiétude voir la révolte de
27 certains citoyens.

29 Gouvernement a décidé de rendre la reconnaissance faciale obligatoire dans certains cas, il
30 s'agit plus précisément du cas de personnes souhaitant acheter un nouveau téléphone portable
31 aujourd'hui tous les Chinois qui veulent acheter un nouveau téléphone portable doivent faire
32 enregistrer leur visage pour pouvoir être reconnu grâce à la reconnaissance faciale la Chine
33 veut en effet pouvoir identifier tous les utilisateurs grâce à leur numéro de téléphone et à leur
34 visage. Alors, il faut savoir que cette obligation de d'enregistrer les citoyens sous leur identité
35 réelle, elle l'est déjà en vigueur en Chine depuis 2013.

37 Elle vise surtout cette initiative à surveiller les réseaux sociaux en effet la majorité des
38 réseaux sociaux en Chine ont l'obligation de connaître l'identité réelle des utilisateurs mais ce
39 qui est nouveau ici, c'est justement le recours à la reconnaissance faciale alors pourquoi est-ce
40 que les citoyens chinois sont particulièrement inquiets de cette mesure et bien surtout parce
41 qu'il craignent que leurs données ainsi enregistré puisse être transmises et revendus à d'autres
42 entreprises en effet, il en effet par rare en Chine de constater des fuites de données
43 personnelles des personnes mal intentionnées qui revendent les données personnelles des
44 utilisateurs à d'autres entreprises.

46 alors en général, il s'agit de donner telles que les adresses ou numéro de téléphone mais dans
47 ce cas-ci et ce qui inquiète plus particulièrement les citoyens, c'est que ça concerne leur

48 donner biométrique, c'est-à-dire leurs empreintes digitales les scans de leur visage ou encore
49 de l'iris de leur œil. Un rapport publié récemment un effet indiqué qu'environ 10% des
50 applications sur téléphone portable était soupçonnées de collecter de manière excessive les
51 données biométriques de leurs utilisateurs.

52
53 récemment encore c'est un autre exemple qui semble témoigner de la méfiance croissante des
54 citoyens chinois. un professeur qui voulait se rendre dans une réserve d'animaux non loin de
55 Shanghai à décider de refuser d'entrer dans la réserve pourquoi parce que les gardiens à
56 l'entrée voulaient scanner son visage. Il l'a refusé que son visage soit scanné et il a ensuite
57 décidé d'attaquer en justice la direction d'une réserve d'animaux ce genre de cas est tout à fait
58 nouveau et sans précédent en Chine et montre donc que les citoyens en ont marre mais qu'est-
59 ce qui pousse le gouvernement chinois à imposer la reconnaissance faciale dans tous les
60 aspects du quotidien. Alors officiellement évidemment comme je l'ai dit, il s'agit d'une mesure
61 visant à assurer la sécurité des citoyens mais il faut savoir également que le gouvernement
62 chinois est désireux de faire de la Chine un leader technologique par conséquent le
63 gouvernement, n'hésite pas à offrir son soutien à des programmes de reconnaissance faciale
64 ou d'intelligence artificielle en général.

65
66 En Chine on compte environ 850 millions d'utilisateurs de téléphones portables, ça constitue
67 une précieuse base de données pour les entreprises et les administrations alors pour tenter de
68 rassurer la population. Le gouvernement a publié des directives. Non contraignantes visant à
69 réglementer la collecte et l'utilisation des données. Mais à l'heure actuelle la Chine ne dispose
70 pas de vraies législations à ce sujet. Alors apparemment une loi est en cours de rédaction mais
71 suffira-t-elle à apaiser les citoyens. les citoyens chinois viaen l'avenir nous le dire. Merci.

T3 – Transkript 3 –Übersetzung des Transkriptes von Google Übersetzer

1 Hallo, ich spreche mit Ihnen über die Gesichtserkennungstechnologie und insbesondere über
2 ihre weit verbreitete Verwendung in China und fange an. Sie haben sicherlich alle von der
3 Gesichtserkennungstechnologie gehört, es ist eine Technologie, mit der Sie eine Person dank
4 ihres Gesichts vollständig automatisch erkennen können. und warum kennen wir sie gut, nur
5 weil diese Technologie jetzt beginnt, in unser tägliches Leben einzugreifen, ihr Ziel ist es, das
6 Leben der Bürger zu erleichtern.

8 Nehmen wir das Beispiel von Mobiltelefonen, die mit Gesichtserkennung entsperrt werden
9 können. Dies ermöglicht es dem Benutzer, sich kein kompliziertes Passwort merken zu
10 müssen und es nicht jedes Mal verschlüsseln zu müssen, wenn er sein Telefon verwenden
11 möchte. Und wenn wir über neue Technologien und ihre Anwendung sprechen, sind wir oft
12 gezwungen, auf China zu verweisen, China ist oft sehr an technologischen Fortschritten
13 interessiert, und das Beispiel der Gesichtserkennung ist da keine Ausnahme.

15 Die Gesichtserkennung ist in der Tat bereits bequem im täglichen Leben der chinesischen
16 Bürger installiert. Es wird zum Beispiel an Supermarktkassen zum Bezahlen von
17 Lebensmitteln verwendet. und die Chinesen im Allgemeinen scheinen der Verwendung dieser
18 Technologie recht wohlgesonnen zu sein, auch wenn sie manchmal etwas aufdringlich wirken
19 kann. Ihrer Meinung nach erleichtert die Gesichtserkennung ihr Leben und ist
20 gleichbedeutend mit vielen zusätzlichen Vorteilen. Ihrer Meinung nach ermöglicht die
21 Gesichtserkennung die Gewährleistung der öffentlichen Sicherheit.

23 Tatsächlich wurde die Gesichtserkennung in China ursprünglich hauptsächlich zur
24 Gewährleistung der Sicherheit eingesetzt. Insbesondere ermöglicht es, die Identitätskontrollen
25 im Verkehr zu beschleunigen, und ermöglicht es auch, gehorsame Bürger zu identifizieren.
26 Eine kürzlich von der chinesischen Regierung ergriffene Maßnahme scheint jedoch zu weit zu
27 gehen, sie hat die Besorgnis und sogar die Revolte einiger Bürger geweckt.

29 Die Regierung hat beschlossen, die Gesichtserkennung in bestimmten Fällen obligatorisch zu
30 machen, insbesondere im Fall von Menschen, die ein neues Mobiltelefon kaufen möchten.
31 Heute müssen alle Chinesen, die ein neues Mobiltelefon kaufen möchten, ihr Gesicht
32 registrieren, um erkannt zu werden zur Gesichtserkennung China will alle Nutzer anhand ihrer
33 Telefonnummer und ihres Gesichts identifizieren können. Sie sollten also wissen, dass diese
34 Verpflichtung, Bürger unter ihrer echten Identität zu registrieren, in China bereits seit 2013 in
35 Kraft ist.

37 Es zielt hauptsächlich auf diese Initiative ab, soziale Netzwerke zu überwachen. Tatsächlich
38 sind die meisten sozialen Netzwerke in China verpflichtet, die wahre Identität der Benutzer zu
39 kennen, aber was hier neu ist, ist genau die Verwendung der Gesichtserkennung. Warum also
40 sind die chinesischen Bürger besonders besorgt darüber? Maßnahme und vor allem weil sie
41 befürchten, dass ihre so erfassten Daten tatsächlich an andere Unternehmen übermittelt und
42 weiterverkauft werden könnten, ist es in China tatsächlich selten, dass personenbezogene
43 Daten von Personen mit bösen Absichten weitergegeben werden, die personenbezogene
44 Daten von Benutzern an andere Unternehmen weiterverkaufen.

46 Im Allgemeinen geht es also darum, Adressen oder Telefonnummern anzugeben, aber in
47 diesem Fall, und was die Bürger besonders beunruhigt, geht es darum, ihnen biometrische

48 Daten zu geben, d. h. ihre Fingerabdrücke, die Scans ihres Gesichts oder der Iris ihres Auges.
49 Ein kürzlich veröffentlichter Bericht weist darauf hin, dass etwa 10 % der
50 Mobiltelefonanwendungen verdächtigt wurden, übermäßig biometrische Daten von ihren
51 Benutzern zu sammeln.
52
53 in jüngster Zeit ist es ein weiteres Beispiel, das das wachsende Misstrauen der chinesischen
54 Bürger zu bezeugen scheint. Ein Professor, der in ein Tierreservat unweit von Shanghai
55 wollte, verweigerte den Zutritt zum Reservat, weil die Wachen am Eingang sein Gesicht
56 scannen wollten. Er weigerte sich, sein Gesicht scannen zu lassen und beschloss dann, die
57 Verwaltung eines Tierreservats zu verklagen. Diese Art von Fall ist in China ziemlich neu
58 und beispiellos und zeigt daher, dass die Bürger in China die Nase voll haben, aber was die
59 chinesische Regierung dazu treibt, Gesichtserkennung in China durchzusetzen alle Aspekte
60 des täglichen Also offiziell offensichtlich, wie ich sagte, dies ist eine Maßnahme, um die
61 Sicherheit der Bürger zu gewährleisten, aber Sie sollten auch wissen, dass die chinesische
62 Regierung daran interessiert ist, China zu einem technologischen Führer zu machen, deshalb
63 die Regierung, n zögern Sie nicht, Gesichtserkennung oder künstliche Intelligenz zu
64 unterstützen Programme im Allgemeinen.
65
66 In China gibt es ungefähr 850 Millionen Mobiltelefonbenutzer, es stellt eine wertvolle
67 Datenbank für Unternehmen und Verwaltungen dar, um zu versuchen, die Bevölkerung zu
68 beruhigen. Die Regierung hat Richtlinien Nicht verbindlich, um die Erhebung und
69 Verwendung von Daten zu regeln. Aber derzeit hat China keine wirkliche Gesetzgebung zu
70 diesem Anscheinend wird also ein Gesetz entworfen, aber wird es ausreichen, um die Bürger
71 zu Chinesische Bürger über die Zukunft erzählen uns. Vielen Dank.

T4 – Transkript 4 – Übersetzung des Transkriptes von Google Übersetzer inklusive der Pausenmarkierungen

1 Hallo, ich spreche mit Ihnen über die Gesichtserkennungstechnologie und insbesondere über
2 ihre weit verbreitete Verwendung in China und fange an. • Sie haben sicherlich alle von der
3 Gesichtserkennungstechnologie gehört, • es ist eine Technologie, mit der Sie eine Person dank
4 ihres Gesichts vollständig automatisch erkennen können. • • und warum kennen wir sie gut, nur
5 weil diese Technologie jetzt beginnt •, in unser tägliches Leben einzugreifen, ihr Ziel ist es, das
6 Leben der Bürger zu erleichtern. • •

7
8 Nehmen wir das Beispiel von Mobiltelefonen, • die mit Gesichtserkennung entsperrt werden
9 können. • Dies ermöglicht es dem Benutzer, • sich kein kompliziertes Passwort merken zu
10 müssen und es nicht jedes Mal verschlüsseln zu müssen, • wenn er sein Telefon verwenden
11 möchte. • • Und wenn wir über neue Technologien und ihre Anwendung sprechen, sind wir oft
12 gezwungen, • auf China zu verweisen, • China ist oft sehr an technologischen Fortschritten
13 interessiert, • und das Beispiel der Gesichtserkennung ist da keine Ausnahme. •

14
15 Die Gesichtserkennung ist in der Tat bereits bequem im täglichen Leben der chinesischen
16 Bürger installiert. • Es wird zum Beispiel an Supermarktkassen zum Bezahlen von
17 Lebensmitteln verwendet. und die Chinesen im Allgemeinen scheinen der Verwendung dieser
18 Technologie recht wohlgesonnen zu sein, • auch wenn sie manchmal etwas aufdringlich wirken
19 kann. • Ihrer Meinung nach erleichtert die Gesichtserkennung ihr Leben und ist gleichbedeutend
20 mit vielen zusätzlichen Vorteilen. • Ihrer Meinung nach ermöglicht die Gesichtserkennung die
21 Gewährleistung der öffentlichen Sicherheit. •

22
23 Tatsächlich wurde die Gesichtserkennung in China ursprünglich hauptsächlich zur
24 Gewährleistung der Sicherheit eingesetzt. • • Insbesondere ermöglicht es, • die
25 Identitätskontrollen im Verkehr zu beschleunigen, und ermöglicht es auch, gehorsame Bürger
26 zu identifizieren. • • Eine kürzlich von der chinesischen Regierung ergriffene Maßnahme
27 scheint jedoch zu weit zu gehen, sie hat die Besorgnis und sogar die Revolte einiger Bürger
28 geweckt. • •

29
30 Die Regierung hat beschlossen, • die Gesichtserkennung in bestimmten Fällen obligatorisch zu
31 machen, • insbesondere im Fall von Menschen, die ein neues Mobiltelefon kaufen möchten. •
32 Heute müssen alle Chinesen, • die ein neues Mobiltelefon kaufen möchten, ihr Gesicht
33 registrieren, • um erkannt zu werden zur Gesichtserkennung China will alle Nutzer anhand ihrer
34 Telefonnummer und ihres Gesichts identifizieren können. • • Sie sollten also wissen, dass diese
35 Verpflichtung, • Bürger unter ihrer echten Identität zu registrieren, in China bereits seit 2013 in
36 Kraft ist. •

37
38 Es zielt hauptsächlich auf diese Initiative ab, • soziale Netzwerke zu überwachen. • • Tatsächlich
39 sind die meisten sozialen Netzwerke in China verpflichtet, • die wahre Identität der Benutzer
40 zu kennen, aber was hier neu ist, ist genau die Verwendung der Gesichtserkennung. • • Warum
41 also sind die chinesischen Bürger besonders besorgt darüber? Maßnahme und vor allem weil
42 sie befürchten, • dass ihre so erfassten Daten tatsächlich an andere Unternehmen übermittelt
43 und weiterverkauft werden könnten, • ist es in China tatsächlich selten, • • dass
44 personenbezogene Daten von Personen mit bösen Absichten weitergegeben werden, die
45 personenbezogene Daten von Benutzern an andere Unternehmen weiterverkaufen.

46
47 Im Allgemeinen geht es also darum, • Adressen oder Telefonnummern anzugeben, aber in
48 diesem Fall, und was die Bürger besonders beunruhigt, • geht es darum, • ihnen biometrische
49 Daten zu geben, d. • h. • ihre Fingerabdrücke, • die Scans ihres Gesichts oder der Iris ihres
50 Auges. Ein kürzlich veröffentlichter Bericht weist darauf hin, • dass etwa 10 % der
51 Mobiltelefonanwendungen verdächtigt wurden, • übermäßig biometrische Daten von ihren
52 Benutzern zu sammeln.
53
54 in jüngster Zeit ist es ein weiteres Beispiel, • das das wachsende Misstrauen der chinesischen
55 Bürger zu bezeugen scheint. • • Ein Professor, der in ein Tierreservat unweit von Shanghai
56 wollte, verweigerte den Zutritt zum Reservat, • weil die Wachen am Eingang sein Gesicht
57 scannen wollten. • Er weigerte sich, • sein Gesicht scannen zu lassen und beschloss dann, die
58 Verwaltung eines Tierreservats zu verklagen. Diese Art von Fall ist in China ziemlich neu und
59 beispiellos und zeigt daher, • dass die Bürger in China die Nase voll haben, aber was die
60 chinesische Regierung dazu treibt, Gesichtserkennung in China durchzusetzen alle Aspekte des
61 täglichen Also offiziell offensichtlich, • • wie ich sagte, dies ist eine Maßnahme, um die
62 Sicherheit der Bürger zu gewährleisten, • aber Sie sollten auch wissen, dass die chinesische
63 Regierung daran interessiert ist, • • China zu einem technologischen Führer zu machen, •
64 deshalb die Regierung, n zögern Sie nicht, Gesichtserkennung oder künstliche Intelligenz zu
65 unterstützen Programme im Allgemeinen. • •
66
67 In China gibt es ungefähr 850 Millionen Mobiltelefonbenutzer, • es stellt eine wertvolle
68 Datenbank für Unternehmen und Verwaltungen dar, • um zu versuchen, • die Bevölkerung zu
69 beruhigen. • • Die Regierung hat Richtlinien Nicht verbindlich, • um die Erhebung und
70 Verwendung von Daten zu regeln. • • Aber derzeit hat China keine wirkliche Gesetzgebung zu
71 diesem Anscheinend wird also ein Gesetz entworfen, aber wird es ausreichen, um die Bürger
72 zu Chinesische Bürger über die Zukunft erzählen uns. • Vielen Dank.

T5 – Transkript 5 – Selbst erstelltes Transkript der Humandolmetschung 1

1 Guten Tag, ich möchte heute über die Technologie der GeSICHTSerkenennung insbesondere
2 der Nutzung in China sprechen. ((ea)) Lassen Sie mich beginnen. ((ea)) Sie haben bestimmt
3 schon alle gehört von dieser Technologie der Gesichtserkennung, die es • ermöglicht, ((ea))
4 automatisch die Gesichter von Personen • ((Zettel)) zu erkennen. Aber warum ist uns diese
5 Technologie überhaupt bekannt? ((ea)) Es ist SO, dass diese Technologie: momentan beginnt,
6 in unserem Alltag omnipräsent zu werden und in unserem Alltag Einzug zu halten. ((ea))
7 Diese Technologie ermöglicht es uns, unser Leben EINFacher zu gestalten; ((ea))
8
9 beispielsweise wenn wir unsere Smartphones via Gesichtserkennung entsperren möchten.
10 ((ea)) Die Nutzer*innen müssen sich dabei nicht an komplizierte Passwörter erinnern oder an
11 ihren Code, wenn sie ihr Telefon benutzen möchten. ((ea)) Aber immer, wenn man über
12 Technologie spricht, muss man auch CHIna erwähnen, ((ea)) weil China • • ((ea)) ein großer
13 Fan von Technologie ist und die Technologie auch immer vorANTreibt. ((ea)) Und die
14 Gesichtserkennung bildet hier keine Ausnahme.
15
16 ((ea)) Schon jetzt ist die Gesichtserkennung ((ea)) im: Alltag Chinas nicht mehr
17 wegzudenken; ((ea)) zum Beispiel an Supermarktkassen kann man seine Einkäufe via
18 Gesichtserkennung beZAHLen. ((ea)) • • Und im Allgemeinen ist diese Technologie auch
19 sehr beliebt bei den Bürgerinnen und Bürgern Chinas; ((ea)) auch wenn sie manchmal ((ea))
20 sehr stark in das Leben: der Menschen EINGreift. ((ea)) Doch es gibt sehr viele Vorteile.
21
22 ((ea)) Die Chinesen denken, • dass es: ihnen das Leben erLEICHTert; zum Beispiel in
23 Hinblick auf die öffentliche Sicherheit. ((ea)) Denn das war der ursprüngliche Zweck, warum
24 diese Gesichtserkennung dort eingeführt wurde. ((ea)) So geh:t es zum Beispiel SCHNELLeR,
25 die Identität zu kontrollieren in öffentlichen Transportmitteln ((ea)) und • Menschen, die sich
26 nicht an die Regeln halten zu identifizieren. • • ((ea)) Es wurde jedoch auch ahm: ein
27 VORSchlag der Regierung Chinas umgesetzt, de:r/ ((ea)) ahm: • • der Proteste in der
28 Bevölkerung hervorgerufen hat;
29
30 ((ea)) und zwar ist es so, dass man sich verPFLICHTend in bestimmten Fällen mit seinem
31 Gesicht registrieren muss; ((ea)) genauer gesagt • müssen sich alle Chinesinnen und Chinesen,
32 ((ea)) ahm wenn sie ein neues Telefon, ein neues Smartphone kaufen möchten ((ea)) mit
33 ihrem Gesicht registrieren. ((schluckt)) Und das Ziel ist es, ((ea)) • • SO alle Nutzerinnen und
34 Nutzer von Telefonen anhand ihres Gesichts identifizieren zu können. ((ea)) Ahm: ((2s)) da:s
35 ahm: Programm gibt es seit zweitausendDREIzehn
36
37 ahm • • ((ea)) u:nd es ist SO, dass zum Beispiel auch auf Social Media ((ea)) die
38 tatSÄCHliche Identität der registrierten Nutzer*innen ((ea)) ahm erfasst und überprüft wird.
39 • • ((ea)) WaRUM sind die Menschen in China so beUNruhigt deswegen? ((ea)) Sie fürchten,
40 dass ihre Daten an andere Unternehmen verKAUFT werden. ((ea)) Das ist in China nicht
41 selten der Fall, dass persönliche Daten an Unternehmen weiterverkauft werden.
42
43 ((ea)) Normalerweise handelt es sich dabei um: • Daten, wie Adressen oder Telefonnummern.
44 ((ea)) Doch in diesem Fall geht es um andere Daten, nämlich die bioMETrischen Daten. ((ea))
45 Das können Fingerabdrücke sein • oder auch die Iris des Auges. • • ((ea)) Es: ist auch bekannt,
46 dass 10% der Apps auf Smartphones exzessiv diese Daten sammeln ((ea)) und: deshalb die
47 Bevölkerung sehr misstrauisch ist.

48
49 ((ea)) Es ist erst vor kurzem der Fall eingetreten, dass ein: • ((ea)) Forscher in einem Reservat
50 in der Nähe von Shanghai keinen Zutritt erhalten hat, ((ea)) weil er sein Gesicht nicht scannen
51 lassen wollte und das ist nur EIN Beispiel; das ist nichts Neues. ((ea)) Und es zeigt, dass die
52 Menschen genug von dieser Technologie haben. ((ea)) Doch: • steht alles unter diesem
53 AsPEKT • ((ea)) der Sicherheit, dass diese • Technologie: • • ((ea)) alle Aspekte des Alltags
54 abdeckt ((ea)) und China hat diesen Ruf ahm:, FÜHREnd auf dem Technologiemarkt zu sein
55 ((ea)) und deshalb sind diese Programme der Gesichtserkennung und der künstlichen
56 Intelligenz insgesamt sehr wichtig für China.
57
58 ((ea)) Und auch weil es ahm: etwa 850 Millionen Nutzer*innen von Smartphones gibt ((ea))
59 und das ist natürlich eine sehr wichtige Datenquelle. ((ea)) U:nd um die Bevölkerung zu
60 beruhigen, ((ea)) ahm: wurde nun angekündigt, die Verwendung der Daten besser zu
61 reguLIERen, aber bis jetzt • ((ea)) gibt es nicht wirklich ((ea)) ahm die Gesetzgebung, die
62 nötig wäre dazu. ((ea)) Und wir werden sehen, wie sich das weiterentwickeln wird in China.
63 ((ea)) Ich • danke Ihnen • recht herzlich für Ihre Aufmerksamkeit.

T6 – Transkript 6 – Selbst erstelltes Transkript der Humandolmetschung 2

1 ((ea)) Heute möchte ich über die technologische Gesichtserkennung sprechen und ((ea)) die
2 vor allem: die verbreitete Anwendung in CHIna. ((ea)) ahm ich beginne. Also • • ((ea)) Sie
3 haben sicher schon mal von der technologischen Gesichtserkennung geHÖRT. ((ea)) Es dient
4 dazu, Personen ((ea)) anhand ihres GeSICHTS • automatisch erkennen zu können ((ea))
5 ((husten)) und • warum spreche ich DARüber? Also heute/ heutzutage ((ea)) ahm mischt sich
6 die: • technologische Gesichtserkennung immer mehr in unser Leben EIN ((ea)) und es soll
7 ((räuspern)) ähm das Leben der Menschen ähm im Grunde LEICHTer machen,
8
9 zum Beispiel ((ea)) • wird die Gesichtserkennung bei HANDys eingesetzt, • wodurch ((ea))
10 die • Benutzer dann keinen • • Pin ahm: mehr eingeben müssen. • • ((ea)) Ahm: wenn man
11 von neuer Technologie und deren Anwendung spricht, dann • spricht man dann oft von China,
12 weil • China sehr verSESSen auf ((ea)) ahm neue Technologie ist u:nd dabei ist auch die •
13 ((ea)) Gesichtserkennung keine • • ((ea)) Ausnahme,
14
15 ((ea)) denn die automatische Gesichtserkennung ist • bereits im Alltag ((ea)) der • Chinesen
16 sehr • • ahm präsent; ((ea)) zum Beispiel wird es an der/ an den Supermarktkassen ((ea)) zum
17 Bezahlen eingesetzt. • • ((ea)) Ahm: ((schluckt)) ähm: grundsätzlich • gilt es als etwas seh:r
18 günstiges, dienliches, das das Leben ((ea)) ahm der Menschen leichter machen soll, • ((ea))
19 obwohl es auch manchmal • sehr aufdringlich sein ((ea)) soll, aber grundsätzlich wird es als
20 etwas gesehen, das viele Vorteile hat ((ea)) u:nd die: • Gesichtserkennung soll • vor allem
21 auch die öffentliche Sicherheit erhöhen.
22
23 • Ahm zu Beginn wurde in China ((ea)) die Gesichtserkennung vor allem eben zur •
24 Gewährleistung der Sicherheit ((ea)) eingesetzt; zum Beispiel in • U-Bahnen ((ea)) dient es
25 der ((räuspern)) Identifikation von Bürgern, ((ea)) die • sich dem Gesetz widersetzen. ((ea))
26 ((räuspern)) Ahm: erst kürzlich ging die Regierung in China jedoch zu WEIT, wo: äh/ was
27 sogar zu Revolten von Bürgern führte,
28
29 • denn die Regierung möchte ((ea)) ah in manchen Bereichen eine verPFLICHTende
30 Gesichtserkennung EINFühren. ((ea)) Zum Beispiel, wenn man sich ein • neues Handy kauft,
31 ((Zettel)) soll man sich verpflichtend/ soll das Gesicht verpflichtend ((ea)) ahm registriert
32 werden, ((ea)) ah:m damit ((räuspern)) die Benutzer ahm mit deren Telefonnummer UND
33 deren Gesicht ((ea)) identifiziert werden können. ((ea)) Ahm ((räuspern)) diese Verpflichtung
34 der Regierung ((ea)) äh: existiert/ ist bereits seit zweitausendDREIzehn in Kraft,
35
36 ahm: • • denn zum Beispiel in den sozialen Netzwerken ((ea)) gibt es diese Verpflichtung
37 bereits, ((ea)) da: ahm müssen alle ah:m/ muss die Identität/ die reelle Identität aller Benutzer
38 bekannt sein ((ea)) und was eben neu ist, ist, dass in diesem Bereich die ((ea)) ahm
39 automatische Gesichtserkennung eingesetzt wird. ((ea)) • • U:nd die Bürger ((räuspern)) sind •
40 • dabei jedoch ahm beUNruhigt, weil ahm die Gesichtserkennung ((ea)) ahm auch dazu
41 benutzt werden kann, dass Daten gesammelt werden • und die dann ahm an andere Firmen
42 weitergegeben • oder verkauft werden können.
43
44 ((ea)) Ahm bei den Daten handelt es sich um Adressen, • Telefonnummern • ((räuspern)) und
45 eben ZUSätzlich ((ea)) um bioMETrische Daten, wie • digitale/ • • den digitalen
46 Fingerabdruck ((ea)) und den Scan der Iris im Aug'. Ahm ((ea)) ah: kürzliche Studien haben

47 auch gezeigt, dass 10% der ahm • Apps bereits • ahm biometrische Daten der Benutzer
 48 speichern
 49
 50 • • ((ea)) u::nd ((2s)) ahm eine letzte/ eine letztes ahm Beispiel, welches • ebenfalls • das
 51 Vertrauen der • ((ea)) ahm chinesischen Bevölkerung sehr • • ((ea)) ahm senkt, ahm ja sehr
 52 verringert hat, ((ea)) war/ ist das Beispiel, wo ein ((ea)) ProFESSor ((ea)) in ein
 53 Wildtierreservat in der Nähe von ShangHAI ((ea)) ahm: sich geweigert HAT, • sein ahm
 54 Gesicht scannen zu lassen und er ist deswegen nicht ((ea)) in das Reservat eingetreten ((ea))
 55 und ist dann äh juristisch gegen das Reservat vorgegangen ((ea)) ahm: • • ((ea)) u:nd das ist
 56 eines der Beispiele, ah: was eben dazu führt, dass die BeVÖLkerung ((ea)) ahm dieses
 57 Thema schon satt hat. ((ea)) U:nd warum will die Regierung die Gesichtserkennung im Alltag
 58 einsetzen? ((ea)) Naja erstens • eben um die Sicherheit • zu garantieren. ((ea)) • • Aber
 59 zweitens auch ((ea)) ahm: wollen sie, • dass China • die technologische Führung • übernimmt
 60
 61 ((ea)) u::nd • • da eine große Anzahl von 850 Millionen ((ea)) ahm: • • Handybenutzern allein
 62 in China ((ea)) • sind ahm, bedeutet das schon eine HOhe Anzahl an Daten. ((ea)) ahm
 63 ((räuspern)) Im Moment hat die Regierung jedoch • äh geäußert, dass sie ((ea)) keine Daten
 64 sammeln und weiter ((ea)) geben werden, • ((ea)) ((räuspern)) jedoch existiert noch KEIN
 65 richtiges • Gesetz dazu; ((ea)) ((husten)) das ist jedoch ((ea)) ahm noch in Planung ((ea)) und
 66 ((Zettel)) wir werden sehen, wie sich das in Zukunft ((husten)) entwickelt.

T7 – Transkript 7 – Selbst erstelltes Transkript der Humandolmetschung 3

1 Guten Tag. • Heute spreche ich über die Technologie der Gesichtserkennung und
2 insbesondere über seine verbreitete Benutzung in China. • Ich beginne jetzt. • Sie haben sicher
3 SCHON mal von der Technologie der Gesichtserkennung gehört. ((ea)) Diese Technologie
4 erlaubt • uns:, ((ea)) Gesichter zu erKENNEN ahm: und zwar • • automatisch. • Warum
5 sprechen wir überhaupt darüber? ((ea)) Eben weil ((2s)) diese Technologie angefangen hat, •
6 einen immer größeren Platz in unserem Leben • zu nehmen ((ea)) und das Ziel daBEI sollte
7 sein, ((ea)) die VerBESSerung oder ErLEICHTerung • verschiedener Prozesse im Alltag zu
8 sein.

9
10 • • ((ea)) Ein Beispiel wäre die Benutzung von Handys. Wer sich jetzt nicht unbedingt
11 wünscht, ((ea)) jedes Mal beim Entsperrten der/ des Handys ein Passwort einzugeben
12 beziehungsweise • dieses Passwort sich zu merken, ((ea)) kann • eine: •
13 Gesichtserkennungstechnologie benutzen. ((ea)) • • Wenn man darüber redet, dann muss man
14 auch ((ea)) über China reden, • denn ((ea)) China ganz viele Fortschritte im technologischen
15 Bereich allgemein, aber eben AUCH im Bereich der Gesichtserkennung, macht.

16
17 ((ea)) • • Gesichtserkennung gehört tatSÄCHlich zu dem ALLtag der chinesischen Bürger.
18 ((ea)) Ein einfaches Beispiel wäre die: äh Kästchen in den Supermark/ Supermärkten an der
19 Kassa, wo:durch man seine Einkäufe machen kann. • • ((ea)) Eigentlich sehen viele • Chineser
20 diese Technologie • in einem positiven Licht, ahm selbst wenn sie vielleicht aufdringlich
21 scheint. ((ea)) Sie meinen, dass es ihnen erlaubt, • etwas ahm/ einen einfacheren Alltag zu
22 haben und dass die Technologie die ÖFFentliche Sicherheit gewährleistet ((ea)). China • •
23 sagt ebenfalls/ Die Regierung sagt ebenfalls, dass die: • Gesichtserkennung zu einer •
24 erhöhten Sicherheit in öffentlichen Räumen • • beitragen sollte.

25
26 ((2s)) Zum Beispiel wird/ wird Gesichtserkennung zurZEIT • für öffentliche Verkehrsmittel
27 äh benutzt, um:: eventuelle probleMATische Bürger • identifizieren zu können. • • Aber • es
28 gibt/ es gab vor kurzem einen Vorschlag, • eine neue Maßnahme diesbezüglich einzuführen
29 ((ea)) und diese Maßnahme bereitet schon für viele Sorgen.

30
31 • • ((ea)) Die chinesische Regierung möchte jetzt/ möchte ab jetzt • es verpflichtend machen,
32 ((ea)) ahm jedes Mal, wenn jemand ein Handy kauft durch einen • Gesichtserkennungsprozess
33 zu gehen. • • China möchte somit die Identität ALLER seiner BÜRger anhand von ihren
34 Telefonnummern und • Gesichtsaufnahmen ((2s)) haben möchten oder aufnehmen. • Man
35 muss/ man müsste aber schon dazu sagen, dass • eine/ dass diese Maßnahme seit
36 zweitausendDREIzehn in Kraft getreten ist;

37
38 zwa:r im Bereich der sozialen Netzwerke. Die chinesische Regierung muss jederzeit wissen,
39 was/ welche • Identität also die echte Identität, die hinter Benutzern • steckt. ((ea)) Warum • •
40 bereitet diese Maßn/ diese neue Maßnahme den chinesischen Bürgern Sorgen? • Eben weil es
41 sich/ ahm: es dazu kommen könnte, dass Daten/ dass die Daten, die dadurch gesammelt
42 werden, ahm ganz einfach verkauft werden. ((ea)) Eigentlich kommt das NICHT so selten in
43 China vor. ((ea)) Es ist bis jetzt immer wieder dazu gekommen, dass ahm Daten einfach
44 leaked worden sind an verschiedene Unternehmen. ((2s)) Ahm das ist also nichts Neues,
45 was aber neu ist, sin/ ist die Art von Daten.

47 Es handelt sich dabei um biometrische Daten, ah weil • die • werden zum Beispiel ahm durch
 48 den digitalen • Fingerabdruck oder durch • das Scannen von • Irisen gemacht. Rund/ es wird
 49 angenommen, dass rund 10% aller Apps ahm da/ • dafür ver/ verdächtigt sind, ah zu viele
 50 biometrische Daten von ihren Benutzern zu sammeln.
 51
 52 • • Vor kurzem gab es noch einen Vorfall diesbezüglich, der ebenfalls die Sorgen der • Bürger
 53 verstärkt hat; nämlich ahm wollte ein chinesischer Professor zu einer Naturreserve in
 54 Shanghai GEHen, zum Besuch; ((ea)) aber beim Eintritt wurde er aufgefordert, • sein Gesicht
 55 zu scannen und das hat er abgelehnt. ((ea)) Er hat also nicht mehr reingehen wollen und hat •
 56 eigentlich auch eine Anklage ((ea)) gegen die Naturreserve • erhoben. Dieser • Vor/ Dieser
 57 Fall hat/ ahm ist eigentlich der erste solcher Art in China. • • Warum möchte aber die
 58 chinesische Regierung Gesichtserkennung wirklich in allen Bereichen einführen? ((ea)) Ahm
 59 der öff/ der offizielle Grund dafür ist die: Sicherheit in öffentlichen Räumen, aber eigentlich
 60 möchte China auch ahm an der Spitze der • technologischen Entwicklung in der Welt sein und
 61 das heißt, ((ea)) sie: nutzen quasi jede Gelegen/ Gelegenheit aus, um DIES zu tun.
 62
 63 ((ea)) Es gibt in China 850 Millio:nen Handybenutzer und DAS sorgt natürlich • für eine
 64 rie:sige ahm Menge an Daten. • • Die: Regierung ((ea)) hat schon ahm einige Regelungen
 65 bezüglich der • • Privatsphäre der Benutzer eingeführt, aber ((ea)) noch nicht/ noch kein
 66 echtes Gesetz; daran wird gerade ((ea)) gearbeitet, aber wird das genug sein? Das werden wir
 67 sehen. ((ea)) Vielen Dank für Ihre Aufmerksamkeit.

T8 – Transkript 8 – Selbst erstelltes Transkript der Humandolmetschung 4

1 ((ea)) Hallo, guten Tag. Heute: werde ich über das Thema Gesichtserkennungstechnologien
2 sprechen und darüber, ((ea)) dass diese Technologie • vor allem in China oder die
3 Verwendung dieser Technologie vor allem in CHIna weitverbreitet ist. ((ea)) Ich würde dann
4 beginnen. ((ea)) • • Sie haben wahrscheinlich bereits von Gesichtserkennungstechnologien
5 geHÖRT. ((ea)) Diese Technologie erlaubt es/ erlaubt es, ((2s)) eine Person automatisch UND
6 • anhand ihrer Gesichtszüge zu erkennen. ((ea)) Warum ist diese Technologie • so bekannt?
7 ((ea)) Diese Technologie ist bereits oder hat begonnen, bereits ((2s)) Teil unseres Alltags zu
8 werden. ((schluckt)) Ziel ist es, • das Leben der Bevölkerung zu erleichtern.
9
10 ((ea)) Ein Beispiel dafür, wie diese Technologie eingesetzt ist/ ((ea)) eingesetzt wird, IST • •
11 beispielsweise die Entsperrung des Handys. ((ea)) Man muss sich keine Passwörter mehr
12 merken, oder aber ah Passwörter eintippen. • • ((ea)) Wenn man von neuen Technologien
13 spricht, bleibt einem gar nichts anderes übrig, als China • als Referenz anzuführen. ((aa))
14 Denn China ist versessen auf neue Technologien ((ea)) und dabei stellt die
15 GeSICHTSerkennungstechnologie keine Ausnahme dar.
16
17 ((Zettel)) • • ((ea)) Die Gesichtserkennungstechnologie ist bereits • ah Teil des • Alltags der
18 chinesischen Bevölkerung. ((ea)) Man kann • mithilfe dieser Technologie im Supermarkt
19 seine Einkäufe zahlen. ((ea)) Diese Technologie scheint ahm: bei der chinesischen
20 Bevölkerung gut anzukommen, ((ea)) obwohl sie • sie auch als ((2s)) in die Privatsphäre
21 eingreifend sehen. • Die Vorteile der Technologie ((ea)) sind, dass sie das LEben erleichtern
22 ((ea)) und die Sicherheit der • Bevölkerung • • sicherstellen.
23
24 ((ea)) Man äh/ Sie steigern auch die Kontrolle in den öffentlichen ahm • Verkehrsmitteln und
25 a:hm sie helfen auch dabei, Kriminelle/ • • bei der Verfolgung von Kriminellen. • • ((ea)) Eine
26 Maßnahme, die die chinesische Regierung eingeführt hat, scheint jedoch • zu weit zu gehen
27 ((ea)) • • und bereitet der chinesischen Bevölkerung große Sorge, • und rief auch eine Revolte
28 hervor.
29
30 ((ea)) In manchen Fällen ist die Gesichtserkennung verpflichtend. • Zum Beispiel bei/beim
31 Kauf eines Telefons. ((ea)) Heute/ Heutzutage muss man • seine Gesichtszüge registrieren
32 ((ea)) • • und auch die Telefonnummer. ((Zettel)) • • Ein weiteres Beispiel ist, dass man seit
33 2013, ((ea)) wenn man sich ein Kont/ äh Konto bei den sozialen Medien erstellt, ((ea)) ah:
34 dass man seine a:hm Idi/ Identität • nachweisen muss.
35
36 • • Das trägt dazu bei, dass die sozialen Medien leichter überwacht werden können • • ((ea))
37 und • soziale Medien sind auch verPFLICHTet, die • Identität ihrer Nutzer zu kennen. ((ea))
38 Neu ist jedoch, dass man: auf die • Gesichtserkennungstechnologie zurückgreift. • • ((ea))
39 Warum ist die chinesische Bevölkerung deshalb besorgt? ((ea)) Sie glauben, dass • ihre Daten
40 übertragen und an Firmen • verkauft werden. ((ea)) Datendiebstahl • kommt häufig vor • und
41 ahm: werden mit bösen Absich/ und persönliche Daten werden mit bösen Absichten an
42 Firmen verkauft.
43
44 ((ea)) Sehr beliebt sind Adressen und • Telefonnummern. • • In diesem Fall ((ea)) sorgen sich/
45 sorgt sich die chinesische Bevölkerung ((ea)) um die biometrischen Daten; ((ea)) • • vor allem
46 geht es um den digitalen Fußabdruck, die Iris und die • Gesichtszüge. ((schluckt)) 10% der

47 auf den Handys installierten Apps stehen unter Verdacht, ((ea)) • • biometrische Daten der
 48 Nutzer in GROSSEM Ausmaß zu sammeln.
 49
 50 ((Zettel)) • • ((ea)) Ein weiteres Beispiel dafür, dass das Misstrauen ((2s)) der Bevölkerung
 51 steigert, ist ahm:, dass ein Professor • Zutritt zu einem Wildreservat haben wollte, ((ea)) aber
 52 der • Wachmann ahm wollte sein • Gesicht einscannen, ((ea)) was der • Professor ihm aber
 53 nicht erlaubt hat. Desha/ Deswegen wurde ihm der Zutritt • in das/ in das Wildreservat nicht
 54 gewährt. ((ea)) Er sah sich dadurch gezwungen, ah: rechtliche Schritte gegen die Direktion
 55 des Wildreservats ((3s)) vorzunehmen. ((ea)) • • Leider gibt es keinerlei • Präzedenzfälle und
 56 dieses Beispiel zeigt auch, dass: ((2s)) Gesichtstechnologien bei der Bevölkerung • großen
 57 UNmut hervorrufen. ((ea)) Warum möchte die Regie/ chinesische Regierung
 58 Gesichtserkennungstechnologie ((ea)) Teil des Alltags machen? ((ea)) • • Einerseits handelt es
 59 sich um eine Maßnahme zur • Gewährleistung der Sicherheit der Bevölkerung. Andererseits
 60 möchte d/ die Regierung zum Leader der Technologiewelt werden • ((ea)) und bietet daher
 61 Firmen, die • • sich mit der/ mit Gesichtserkennungsprogrammen und Artificial
 62 InTELLigence beschäftigen, ((ea)) ahm: Unterstützung an.
 63
 64 Es gibt etwa 850 Millionen Menschen, die Handys nutzen. ((ea)) Dies entspricht • ahm: einer
 65 sehr • • wertvollen • Datenbasis für • • Unternehmen und • • die Regierung. Es wurde ein
 66 Versuch zur Beruhigung der Bevölkerung unternommen ((ea)) und zwar hat die Regierung
 67 eine RICHTlinie veröffentlicht; ((ea)) zur • • Sammlung und Verwendung von persönlichen
 68 Daten. ((ea)) Diese ist jedoch UNverbindlich. ((ea)) Es gibt leider keine/ kein Gesetz in dieser
 69 Hinsicht, • obwohl anscheinend gerade eines in Arbeit ist. • Es stellt sich allerdings die Frage:
 70 "Reicht das?". ((ea)) Die Zukunft wird es wohl zeigen.

Abstract (Deutsch)

Ziel dieser Masterarbeit war es herauszufinden, welche Möglichkeiten und Grenzen sich bei der Verwendung von ‚Google Übersetzer‘ als Software für maschinelles Dolmetschen einer Rede im Modus des Konsekutivdolmetschens im Vergleich zu Humandolmetscher:innen ergeben. Dabei wurden die Dolmetschleistungen von vier Dolmetscher:innen mit jener von Google Übersetzer anhand eines Bewertungsschemas, welches durch Auseinandersetzung mit der Literatur erarbeitet wurde, beurteilt. Die Ergebnisse zeigen, dass sich beim maschinellen Dolmetschen mit Google Übersetzer auf verschiedenen Ebenen Probleme ergeben. So enthält der Zieltext gegensätzliche Aussagen im Vergleich zum Ausgangstext sowie einige unverständliche Passagen. Außerdem müssen bei der Realisierung der Dolmetschung hinsichtlich des Klangs und der Intonation der künstlich erzeugten Stimme einige Abstriche hinsichtlich der Qualität gemacht werden. Andererseits werden auch bei den Humandolmetscher:innen Fehler sichtbar, die allerdings weitaus weniger schwerwiegend für eine Gesamtbewertung sind. Dementsprechend zeigt die Analyse, dass die maschinelle Dolmetschung durch Google Übersetzer abschnittsweise Ergebnisse erzielen kann, die mit jenen von professionellen menschlichen Dolmetscher:innen mithalten können, sich aber noch nicht für den Einsatz bei komplexen und sensiblen Themen eignet.

Abstract (English)

The aim of this master's thesis was to find out which possibilities and limitations arise when using Google Translate as software for machine interpreting of a speech in consecutive interpreting mode in comparison to human interpreters. The interpreting performance of four interpreters was compared with that of Google Translate by using an evaluation scheme that was developed by analysing the literature. The results show that there are problems with machine interpreting with Google Translate on various levels. The target text contains contradictory statements compared to the source text as well as some incomprehensible passages. In addition, some compromises must be made regarding the quality of the sound and intonation of the artificially generated voice in the realisation of the interpretation. On the other hand, errors are also visible in the human interpretations, but these are far less serious for an overall assessment. Accordingly, the analysis shows that machine interpreting by Google Translate can achieve results in some sections that can keep up with those of professional human interpreters, but that this is currently not yet suitable for use with complex and sensitive topics.