



universität
wien

MASTERARBEIT / MASTER'S THESIS

Titel der Masterarbeit / Title of the Master's Thesis

„Der Einfluss von SDH-Untertiteln
auf die filmische Narration“

verfasst von / submitted by

Vanessa Kerstin Fellner, BA BA

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

Wien, 2023 / Vienna 2023

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

UA 070 331 342

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

Masterstudium Translation Deutsch Englisch

Betreut von / Supervisor:

Univ.-Prof. Mag. Dr. Klaus Kaindl

Danksagungen

An dieser Stelle möchte ich bei meiner Familie, vor allem bei meinen Eltern, und bei meinen Freund:innen bedanken, die mich auf meinem Weg begleitet haben und mir mit Rat und Tat zur Seite standen. Und ich möchte mich insbesondere bei meinem Partner Kieran bedanken, der mir im letzten Jahr mit Abstand die größte Stütze war und immer an mich geglaubt hat.

Danke, dass ihr mich unterstützt habt, wenn ich euch gebraucht habe – ich schätze es wirklich sehr, Menschen wie euch an meiner Seite zu haben.

Meine besondere Dankbarkeit gilt meinem Betreuer Univ.-Prof. Mag. Dr. Klaus Kaindl für die fachliche Unterstützung im Masterarbeitsprozess, das direkte und hilfreiche Feedback zu meinen Kapiteln und die vielen Literaturvorschläge.

Ein großes Dankeschön geht auch an alle, die mich während des Masterarbeitsprozesses unterstützt haben – ob es nun zahlreiche Bibliotheks-Schreibtreffen, Korrekturlesungen meiner Kapitel oder einfach nur Unterhaltungen zum Thema waren.

Selbstständigkeitserklärung

Ich versichere, die vorliegende Arbeit selbstständig verfasst zu haben. Ich habe keine anderen als die angegebenen Quellen und Hilfsmittel benutzt. Alle von mir für direkte und indirekte Zitate benutzten Quellen sind nach den Regeln des wissenschaftlichen Zitierens angegeben. Mir ist bekannt, dass beim Verstoß gegen diese Regeln eine positive Beurteilung der Arbeit nicht möglich ist. Ich habe die Arbeit bzw. Teile davon weder im In- noch im Ausland zur Begutachtung als Prüfungsarbeit vorgelegt.

Wien, Mai 2023

Vanessa Fellner, BA BA

Inhalt

Abbildungsverzeichnis	VII
Einleitung	1
1. Multimodalität	4
1.1 Grundlagen und Grundbegriffe	4
1.1.1 Modus – Submodus – Medienvariante	5
1.1.1.1 Bild	7
1.1.1.2 Sprache	8
1.1.1.3 Ton.....	9
1.1.1.4 Musik.....	9
1.1.2 Verarbeitung von Multimodalität	11
1.2 Multimodalität in der Translationswissenschaft.....	12
1.2.1 Der Weg zur Multimodalität	13
1.2.2 Zugrundeliegende Theorie: (Multimodale) Sozialesemiotik	15
1.2.3 Multimodalitätsforschung heute	16
1.2.4 Die Analyse multimodaler Texte.....	18
2. Barrierefreiheit und Translationswissenschaft	21
2.1 Grundlagen der Barrierefreien Kommunikation	21
2.1.1 Definition von Barrierefreiheit und Barrierefreier Kommunikation	21
2.1.2 Teilhabe und Ableismus	22
2.1.3 Kommunikationsbarrieren	24
2.1.4 Arten von Kommunikationseinschränkungen	26
2.1.4.1 Perzeptionseinschränkungen	26
2.1.4.2 Verständniseinschränkungen.....	28
2.1.5 Strategien zum Barriereabbau	29
2.2 Das Verhältnis von Barrierefreiheit und Translation	30
2.2.1 Media Accessibility	32

2.2.2 Anwendungsbereiche der Media Accessibility	33
2.2.3 Wissenschaftliche Verortung.....	35
2.2.4 Access to content.....	37
2.2.5 Access to creation: Accessible Filmmaking.....	37
3. Audiovisuelle Translation und (SDH-)Untertitelung im Wandel der Zeit.....	41
3.1 Audiovisuelle Translation	41
3.1.1 Historische Entwicklung der AVT	42
3.1.2 Untertitelungs- vs. Synchronisationsländer.....	43
3.1.3 Eingang in die Translationswissenschaft.....	44
3.1.4 Terminologische Diskussion	45
3.2 (SDH-)Untertitelung.....	46
3.2.1 Definitorische Annäherung	46
3.2.2 Besonderheiten der SDH-Untertitelung	48
3.2.3 Die Heterogenität der Zielgruppe.....	49
3.2.4 Überblick über die historische Entwicklung der SDH-Untertitelung.....	51
3.2.5 SDH-Richtlinien.....	56
3.3 Aktuelle Forschungsschwerpunkte.....	60
3.3.1 Lesbarkeit, Lesegeschwindigkeit und Verständlichkeit.....	60
3.3.2 Informationsdichte und Redundanz.....	61
3.3.2 Integration von Symbolen	62
3.3.3 Automatisierungsprozesse.....	63
3.3.4 Kinder als Zielgruppe	64
3.3.5 Sprachenlerner:innen als Zielgruppe.....	65
3.3.6 Narration in audiovisuellen Produkten.....	66
4. Methode und Untersuchungsgegenstand.....	69
4.1 Untersuchungsgegenstand: „Behind Her Eyes“ (2021)	69
4.2 Methodik.....	70

4.3 Ziel der Analyse	73
5. Untersuchung der Situations- und Kontextangaben in „Behind Her Eyes“	75
5.1 Musik.....	75
5.1.1 Musikbeschreibungen.....	75
5.1.1.1 Titelmelodie.....	75
5.1.1.2 Generische Musikbeschreibungen.....	76
5.1.2 Musiktitel und -texte.....	78
5.2 Geräusche	81
5.3 Charaktere.....	83
5.3.1 Akzente.....	83
5.3.1.1 Schottischer Akzent.....	83
5.3.1.2 Französischer Akzent	85
5.3.2 Sprecher:innen-Identifikation.....	87
5.3.3 Andere paraverbale Merkmale	88
5.4 Diskussion der Ergebnisse.....	90
6. Conclusio.....	95
Primärquellen	98
Bibliografie.....	98
Abstracts.....	108

Abbildungsverzeichnis

Abb. 1: Hauptmodi audiovisueller Texte und deren Medienvarianten (Pérez-González 2015: 194, adaptiert nach Stöckl 2004).....	6
--	---

Einleitung

„Accessibility is a form of translation and translation is a form of accessibility, uniting all population groups and ensuring that cultural events, in the broadest sense of the word, can be enjoyed by all.“ (Díaz Cintas et al. 2007: 14)

In vielen Bereichen des alltäglichen Lebens ist das Bewusstsein für die Notwendigkeit von Barrierefreiheit in den letzten Jahren deutlich gestiegen. Dabei bezieht sich die Barrierefreiheit nicht nur auf physische Barrieren wie Treppen oder Toiletten, sondern auch auf andere Bereiche, zu denen auch die Zugänglichkeit von Medien, die *Media Accessibility*, gehört. Diese Form der Barrierefreiheit kann sich auf verschiedene Arten manifestieren, etwa durch Leichte Sprache, die Bereitstellung von Gebärdensprachdolmetscher:innen oder auch das Verbessern der Usability einer Homepage. Eine vielgenutzte Möglichkeit der Barrierefreiheit in Bezug auf audiovisuelle Medien sind Untertitel für Gehörlose und Hörgeschädigte, welche immer häufiger zur Verfügung gestellt werden – ob im Fernsehen oder von Streamingdiensten und sozialen Netzwerken. So verpflichtete sich zum Beispiel der österreichische Fernsehsender ORF dazu, bis zum Jahr 2030 sein gesamtes Angebot auch barrierefrei zugänglich zu machen (vgl. Salzburger Nachrichten 2020). Laut eigener Angabe sind derzeit etwa 75 Prozent aller Programme des Senders barrierefrei via Teletext zuschaltbar (vgl. ORF 2023).

Es gibt trotz des Bestrebens, audiovisuelle Produkte für alle Menschen zugänglich zu machen, bis heute keine einheitlichen SDH-Untertitelungsrichtlinien – weder national noch international –, was unter anderem dazu führt, dass zwischen den einzelnen Sprachen oder auch zwischen Medienhäusern innerhalb einer einzelnen Sprachregion große Unterschiede in der Gestaltung der Untertitel vorkommen können. Diese können sich zum Beispiel auf die farbliche Gestaltung, auf die Kennzeichnung von Situations- und Kontextangaben oder auf Lesegeschwindigkeiten beziehen. Die vorliegende Masterarbeit soll die Situations- und Kontextangaben der SDH-Untertitelungen in der Serie „Behind Her Eyes“, die im Jahr 2021 bei Netflix erschienen ist, analysieren und anschließend feststellen, auf welche Weise diese Angaben Einfluss auf die filmische Narration der Serie haben. Die Forschungsfrage lautet demnach:

„Inwiefern nehmen die Unterschiede in den Situations- und Kontextangaben der SDH-Untertitel in der Netflix-Serie ‚Behind Her Eyes‘ Einfluss auf die filmische Narration?“

Um die Forschungsfrage zu beantworten, sollen die Situations- und Kontextangaben, zum Beispiel Speaker-IDs oder auch Musikbeschreibungen, der deutschen, englischen und französischen Sprachfassung der SDH-Untertitel der genannten Serie anhand eines deskriptiven Übersetzungsvergleichs nach Chaumes Modell zur Analyse von audiovisuellen Texten (2004) untersucht werden. Das Modell vereint sowohl film- als auch translationswissenschaftliche Aspekte audiovisueller Texte und ist daher sehr gut für diese Art der Analyse geeignet.

Die vorliegende Masterarbeit ist in sechs Unterkapitel unterteilt. Da audiovisuelle Produkte aus mehreren bedeutungsgebenden Codes (Bild, Musik, Ton, Sprache) bestehen, werden sie auch als multimodale Texte bezeichnet. Zu Beginn erfolgt daher eine Veranschaulichung der zugrundeliegenden Theorie der multimodalen Kommunikation. Zunächst wird die Multimodalitätstheorie nach Stöckl (2004), die den Rahmen für die spätere Analyse darstellt, detailliert beschrieben, und anschließend wird die Rolle der Multimodalität in der Translationswissenschaft und der heutige Forschungsstand diskutiert.

Untertitel für Gehörlose sind außerdem eine Form der Barrierefreien Kommunikation, weshalb diese Thematik im zweiten Kapitel diskutiert wird. Ähnlich wie im ersten Kapitel werden hier zunächst Grundlagen der Barrierefreiheit beschrieben, bevor im nächsten Schritt besonderes Augenmerk auf das Verhältnis zwischen Translation und Barrierefreiheit gelegt wird. In diesem Teil der Arbeit wird auch die Media Accessibility inklusive ihrer Anwendungsfelder und wissenschaftlichen Verortung genauer beschrieben.

Im dritten Kapitel liegt der Fokus zunächst auf der audiovisuellen Translation, ihrer Entwicklung und ihrem Eingang in die Translationswissenschaft. Darauf folgt eine detaillierte Beschreibung des Untersuchungsschwerpunktes der vorliegenden Masterarbeit, nämlich der Untertitelung für Hörgeschädigte, auch SDH genannt. In diesem Kapitel werden zunächst die Hauptmerkmale von SDH beschrieben; anschließend folgt eine Veranschaulichung der translationswissenschaftlichen Forschungsschwerpunkte in Hinblick auf SDH-Untertitelung. Am Ende dieses Kapitels werden jene Herausforderungen näher beschrieben, die die Multimodalität mit sich bringt und die für die folgende Analyse äußerst relevant sind, nämlich die Narration in audiovisuellen Produkten und Informationsdichte und Redundanz in SDH-Untertitelungen.

Im vierten Kapitel werden die Methode – das Analysemodell für audiovisuelle Texte nach Chaume (2004) – und der Untersuchungsgegenstand (die Netflix-Serie „Behind Her Eyes“) genauer beschrieben.

Anschließend folgt die Untersuchung ausgewählter Beispiele aus der Serie, wobei das Hauptaugenmerk auf den Situations- und Kontextangaben liegt. Im ersten Teil der Analyse

werden jene Merkmale untersucht, die sich auf Musikbeschreibungen beziehen (z.B. allgemeine Musikbeschreibungen, Musiktitel und -texte). Danach liegt der Fokus auf Geräuschbeschreibungen. Im letzten Teil der Untersuchung werden die Untertitel in Hinblick auf jene Angaben untersucht, die sich auf Charaktere beziehen, wie Akzente oder auch Speaker-IDs. Die Ergebnisse werden im Anschluss diskutiert und mit Erkenntnissen aus bereits erfolgten Forschungsarbeiten verglichen.

Im letzten Schritt werden die Erkenntnisse aus dem theoretischen Teil der Masterarbeit sowie die Ergebnisse der durchgeführten Analyse noch einmal zusammengefasst; außerdem erfolgt ein Ausblick für mögliche zukünftige Forschungsarbeiten.

1. Multimodalität

1.1 Grundlagen und Grundbegriffe

Sowohl in der Wissenschaft als auch in der Allgemeingesellschaft galt Translation lange Zeit als rein monomodales Feld, nämlich als Übersetzung von gedruckten Texten. Durch die technischen Entwicklungen der letzten Jahrzehnte jedoch betrifft Translation immer mehr auch multimodale Texte, bei denen etwa geschriebener Text mit Bildern oder Musik verbunden ist. Gambier (2006: 6) argumentiert sogar, dass Bedeutung immer multisemiotisch ist, da Multimodalität eben nicht nur – wie traditionell angenommen – visuelle, akustische und linguistische Elemente vereint, sondern auch Texte mit scheinbar „rein“ linguistischen Merkmalen durch Seitenlayout und Typographie als multimodal bezeichnet werden können. Die Verwobenheit dieser Modi muss von professionellen Sprachdienstleister:innen, etwa technischen Übersetzer:innen, Copywriter:innen, Lokalisierungsexpert:innen und Untertitler:innen, in den Translationsprozess miteinbezogen werden (vgl. O’Sullivan 2013: 2).

Unter *Multimodalität* versteht man im Allgemeinen die Kombination verschiedener bedeutungsgebender *Modi* – etwa Sprache, Bild, Gestik und Mimik – zu einem großen kommunikativen Ganzen. Kress & van Leeuwen, welche das Konzept als erste in das Feld der Translationswissenschaft eingebettet haben, definieren den Begriff als „the use of several semiotic modes in the design of a semiotic product or event, together with the particular way in which these modes are combined“ (Kress & van Leeuwen 2001: 20). Das Konzept existiert seit der Mitte der 1990er-Jahre und hat seitdem in viele akademische Felder, etwa Semiotik, Linguistik, Medienwissenschaften oder auch Soziologie, Eingang gefunden (vgl. Jewitt et al. 2016: 1). Während sich die eben genannten Felder ursprünglich vor allem mit jenen semiotischen Ressourcen beschäftigten, die unter ihren jeweiligen „Zuständigkeitsbereich“ fielen, befasst sich die Multimodalität mit der Synergie semiotischer Ressourcen. Multimodale Produkte jeglicher Art basieren auf dem Zusammenspiel der relevanten gleichzeitig auftretenden Modi (vgl. Pérez-González 2019: 346). Aus diesem Grund hat der Translationsprozess auch dann Auswirkungen auf den gesamten multimodalen Text, wenn nur einer seiner Teile übersetzt wird (vgl. Weissbrod & Kohn 2019: 4).

Baldry & Thibault (2006: 18) definieren audiovisuelle Texte als „composite products of the combined effect of all the resources used to create and interpret them“. Da sie gleichzeitig verschiedene Zeichenrepertoires bedienen, darunter geschriebene und gesprochene Sprache, Musik oder Perspektive, werden sie auch als multimodale Texte bezeichnet. Für diese Vielfalt

an semiotischen Ressourcen stehen dem Publikum auch verschiedene *Medien* zur Verfügung. Medien werden definiert als „the material resources used in the production of semiotic products and events, including both the tools and the materials used“ (Kress & van Leeuwen 2001: 22). Beispiele für Medien sind Bildschirme und Lautsprecher, aber auch Papier oder Software (vgl. Pérez-González 2019: 347). Audiovisuelle Texte sind daher nicht nur multimodal, sondern auch *multimedial*¹, da ein Zusammenspiel an semiotischen Modes mithilfe verschiedener kombinierter Medialitäten produziert und interpretiert wird (vgl. Pérez-González 2015: 186f.). Die Besonderheit dieses Texttyps besteht darin, dass das Verknüpfen der verschiedenen Kommunikationsmodi meist ohne aktive kognitive Anstrengung geschieht. Alle gleichzeitig auftretenden Modi werden zu einer „single unified gestalt in perception“ (Stöckl 2004: 16), welche wir als Rezipient:innen durch unbewusste und routinierte intermodale Verknüpfungen entschlüsseln (vgl. Stöckl 2004: 16). Multimodalität ist, so Stöckl (2004), ein vernetztes System von Entscheidungen – bei der Erstellung eines audiovisuellen Texts entscheidet sich der/die Ersteller:in für jene Modi, dessen Bedeutungspotenzial den Zweck der Kommunikation am besten erfüllen kann. Da alle Modi wiederum eigene Zeichensysteme darstellen, eröffnen sie den Kommunikator:innen ein neues Set an Entscheidungsmöglichkeiten (vgl. Pérez-González 2015: 191).

Kress & van Leeuwen (2001) betonen, dass die Aspekte Modus und Medium nicht immer eindeutig voneinander unterschieden werden können: Einerseits werden Modi immer in medialen Kontexten realisiert und beziehen sich daher immer auf ein Medium, andererseits begünstigen bestimmte Medien auch die Verwendung bestimmter Modi (vgl. Kaindl 2013: 258f.). Sowohl Modus als auch Medium haben also Einfluss auf den Translationsprozess und das Produkt (vgl. Kaindl 2020: 55). Da Translation nicht nur einen Wechsel des Modus, sondern auch des Mediums bedeuten kann, ist es für eine multimodale Kommunikationstheorie unabdingbar, so etwa Kaindl (vgl. 2013: 259) und Stöckl (vgl. 2004: 11), diese so eng verbundenen Aspekte getrennt voneinander zu betrachten.

1.1.1 Modus – Submodus – Medienvariante

Stöckl (2004) unterscheidet in seiner Theorie vier wichtige, teilweise aneinander gekoppelte Charakteristika multimodaler Texte: Hauptmodus, Submodus, Medium und Medienvariante.

¹ Multimodale Texte können sowohl multi- als auch monomedial sein: Audiovisuelles Material, das sowohl über den auditiven als auch den visuellen Kanal empfangen wird, ist *multimedial*, während etwa Radioprogramme oder geschriebene Texte nur entweder gehört oder gesehen werden können und daher als *monomedial* bezeichnet werden (vgl. Pérez-González 2015: 187).

Unter dem Begriff *Hauptmodus* (*core mode*) versteht man dabei die visuellen und auditiven semiotischen Ressourcen, die für das Erstellen und Interpretieren audiovisueller Texte nötig sind, und zwar Ton (sound), Musik (music), Bild (image) und Sprache (language). Diese Hauptmodi „contain sign-repertoires that are deeply entrenched in people’s popular perception of codes and communication“ (Stöckl 2004: 14). Es sind also auch jene Sets an bedeutungsgebenden Ressourcen, auf die wir zurückgreifen, um Meinungen über audiovisuelle Texte zu formulieren, so etwa auch bei Filmkritiken (vgl. Pérez-González 2015: 192).

Hauptmodi sind abstrakte Formen semiotischer Ressourcen, die in bestimmten *Medienvarianten* (*medial variants*) ausgedrückt werden müssen, um aufgenommen und verstanden werden zu können (vgl. Stöckl 2004: 14). In der untenstehenden Abbildung wird dieses Netzwerk an Hauptmodi und Medienvarianten bildlich dargestellt.

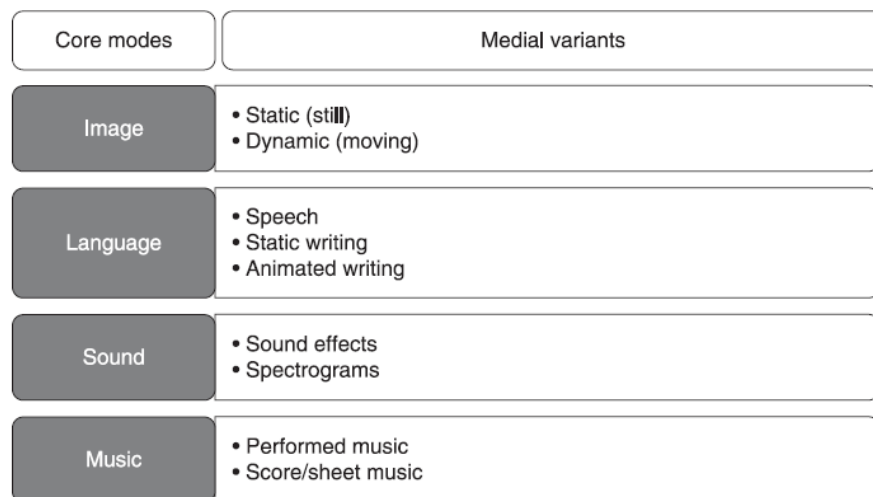


Abb. 1: Hauptmodi audiovisueller Texte und deren Medienvarianten (Pérez-González 2015: 194, adaptiert nach Stöckl 2004)

Die Modi Ton und Musik können sowohl durch auditive als auch durch visuelle Medien ausgedrückt werden. Neben der zweidimensionalen Repräsentation mithilfe von Elektrospektrogrammen kann der Ton eines audiovisuellen Textes auch durch aufgenommene Sprache, also Dialog oder Soundeffekte, realisiert werden. Genauso vollzieht sich auch die Darstellung des Mediums Musik – während es meistens als Teil des Soundtracks realisiert wird, sind beispielsweise auch ausgedruckte Musiknoten eine Ausdruckweise dieses Mediums (vgl. Pérez-González 2015: 194).

Auch Bild und Sprache können verschiedenartig ausgedrückt werden. Der Modus Bild kann sowohl statisch, beispielsweise durch Verwendung von Freeze-Frames, als auch dynamisch, also durch bewegte Bilder, realisiert werden. Die Sprache in audiovisuellen Medien wird zumeist statisch durch Unter- oder Zwischentitel dargestellt, etwa um die Geschehnisse

im Film zu kommentieren oder den gesprochenen Dialog zu übersetzen. Immer öfter jedoch machen gerade Fansubber:innen² Gebrauch von dynamischen Darstellungsmöglichkeiten von Sprache, etwa mithilfe von Untertiteln, die wie beim Karaoke den gerade gesungenen bzw. gesprochenen Teil hervorheben. Diese Tatsache lässt sich damit begründen, dass Fansubber:innen oft vor allem im Bereich der Anime-Übersetzung tätig sind und die Rezipient:innen dieser Übersetzungen performative und immersive Erlebnisse wertschätzen. Durch die Karaoke-artigen Untertitel können diese das Geschehen am Bildschirm noch besser verstehen und fühlen sich miteinbezogen. Auch Regisseur:innen verwenden dynamische Untertitel als Stilmittel in ihren Filmen. Ein Beispiel hierfür ist Timur Bekmambetovs Film „Night Watch“ (2005): Während der Protagonist im Wasser schwimmt, hört er eine Stimme, die ihn wiederholt zu sich bittet („Come to me“, „Come, I’m waiting“), woraufhin seine Nase anfängt zu bluten. Ähnlich wie das Blut des Protagonisten im Wasser verschwimmen auch die Untertitel, die diese Stimme darstellen sollen (vgl. Pérez-González 2015: 195ff.).

Neben den spezifischen Medienvarianten verlangt jeder Hauptmodus außerdem ein Set an *Submodi* (*sub-modes*). Wenn ein:e Kommunikator:in sich für einen bestimmten Hauptmodus entschieden hat, stehen ihr:ihm eine Anzahl an Submodi zur Verfügung, um diesen Hauptmodus in einem audiovisuellen Text zu realisieren. Das Zusammenspiel der Submodi bestimmt darüber hinaus über (die Bedeutung und) den kommunikativen Wert des Hauptmodus, da dieser mehrere zugehörige Sets an Submodi haben kann. Die konkrete Darstellung eines Modus in Text und Diskurs ist, so Stöckl (vgl. 2004: 15), dabei immer mehr als die Summe seiner Einzelteile. Im Anschluss folgt eine Skizzierung jener Haupt- und Submodi, die im weiteren Verlauf der Arbeit, also für die Analyse von SDH-Untertitelung, relevant sind.

1.1.1.1 Bild

Der Modus Bild besitzt im Vergleich zu den anderen Hauptmodi die meisten Submodi, da die meisten von ihnen sowohl statisch als auch dynamisch realisiert werden können. Für die SDH-Untertitelung sind diese jedoch weniger relevant, da dieser Modus von hörgeschädigten Rezipient:innen vollständig empfangen werden kann. Submodi des Modus Bild sind unter anderem Elemente, Vektoren, Farbe, Belichtung, Kameraeinstellung, Kamerabewegungen und visuelle Effekte.

² Als *Fansubbing* bezeichnet man das Untertiteln von Filmen und Serien durch die Fan-Community (vgl. Pérez-González 2015: 17).

Elemente – dazu zählen etwa Menschen, Objekte oder auch abstrakte visuelle Darstellungen – sind die kleinsten visuellen Bausteine dieses Hauptmodus. In audiovisuellen Texten werden visuelle und verbale Elemente durch den „language-image-link“ (Stöckl 2004: 21) verbunden: Die einzelnen Elemente verschmelzen miteinander, formen intermodale Bedeutungsbeziehungen und helfen den Rezipient:innen, den multimodalen Text zu verstehen (vgl. Pérez-Gonzales 2015: 214).

1.1.1.2 Sprache

Der Modus Sprache wird hauptsächlich mithilfe von Systemen verbaler Bedeutungsträger, darunter Grapheme und Phoneme sowie lexikalische, grammatische und syntaktische Bausteine, realisiert. Diese Submodi nehmen Einfluss auf die Realisierung von Sprache, unabhängig von ihrer Medienvariante. Sprache kann zwei Ausprägungsarten annehmen: visuell (statische und dynamische Schrift) oder akustisch (paraverbale Charakteristika) (vgl. Stöckl 2004: 13).

Besonders interessant ist, dass die paraverbalen Charakteristika lange Zeit wissenschaftlich unterrepräsentiert waren, obwohl sie, wie Untersuchungen zeigen, großen Einfluss auf die Rezeption multimodaler Produkte haben. So zeigt Bosseaux (2008), dass die Charakterisierung der Hauptperson Buffy aus dem US-amerikanischen Film „Buffy – Der Vampirkiller“ (1992) in der englischen und französischen Version grundlegende Unterschiede aufweist. Buffys Prosodie im englischen Original ist eher monoton, was mit ihrem sonst jugendhaften Auftreten in Gegensatz tritt, während die französische Sprachfassung ein höheres Register anwendet, was Buffy als affektiert darstellen könnte. Ein weiteres Beispiel für paraverbale Charakteristika sind Dialekte. Queen (2004) analysierte die Translation von African American English (AAE) ins Deutsche und fand heraus, dass in den deutschen Synchronfassungen meistens Sprachvariationen verwendet wurden, die mit der deutschen Arbeiterklasse oder mit Jugendsprache assoziiert werden. Diese Lösung wurde wahrscheinlich ausgewählt, um den „clear proletarian overtone“ (Berthele 2000: 607) des Originals beizubehalten (vgl. Queen 2004: 523). Darüber hinaus schien es einen Zusammenhang mit soziodemographischen Faktoren zu geben: Wenn die Sprecher:innen nicht „young, male, urban or involved with street culture“ (Pérez-González 2015: 202) waren, wurden diese paraverbalen Marker meistens nicht eingesetzt. Anders als bei der Synchronisierung werden bei interlingualen Untertiteln meist keine paraverbalen Charakteristika dargestellt. Im Gegensatz dazu sind sie bei der (intralingualen) Untertitelung für Hörgeschädigte ein essenzieller

Bestandteil und werden zumeist in eckigen Klammern dem Gesagten vorangestellt (z.B. [schnell:], [hoch]), sodass hörgeschädigte Zuseher:innen den audiovisuellen Text in gleichem Maße interpretieren können wie das hörende Publikum (vgl. Pérez-González 2015: 201ff.).

1.1.1.3 Ton

Auch Ton kann auf mehrere Arten dargestellt werden: akustisch (als Soundeffekt) oder visuell (als Spektrogramm) (vgl. Stöckl 2004: 13). Da die visuelle Darstellung nur sehr selten in audiovisuellen Texten vorkommt, wird dieser Abschnitt ausschließlich von Soundeffekten handeln. Deren diegetische Submodi, zu denen Intensität, Lautstärke und Qualität zählen (vgl. Stöckl 2004: 13), gehören im Normalfall nicht zum Aufgabengebiet audiovisueller Translator:innen, obwohl sie sie vor einige Herausforderungen stellen können. Da sie ein wichtiger Bestandteil des Interpretationsprozesses audiovisueller Produkte sind, müssen Translator:innen sich ihrer bedeutungsvermittelnden Funktion bewusst sein. Soundeffekte können dabei auf die Übersetzung Einfluss nehmen, zum Beispiel, wenn ein Soundeffekt gleichzeitig den Dialog beeinflusst. Chaume Varela (2004) analysiert etwa einen Fall, in dem ein Charakter eines Films die englische Redewendung „clean as a whistle“ ausspricht, während ein pfeifender Ton zu hören ist. In vielen Sprachen gibt es jedoch keine entsprechende Redewendung, weshalb die Translator:innen entweder auf andere Redewendungen zurückgreifen müssen oder das Produktionsstudio, das auch für die Synchronisation zuständig ist, diesen Ton gänzlich aus der Tonspur löscht. In jedem Fall sind aber besondere Entscheidungsprozesse nötig (vgl. Pérez-González 2015: 206f.).

1.1.1.4 Musik

Die akustische Darstellungsform des Modus Musik – *performed or incidental music* – wird durch eine Vielzahl an Submodi charakterisiert, darunter Melodie, Rhythmus, Geschwindigkeit oder auch den Liedtext. *Performed* bedeutet hierbei, dass die Musik als dramaturgisches Mittel eingesetzt wird, während *incidental music* eine Ergänzung zu dem in der Szene Gesagten darstellt. Visuell kann Musik als einzelne Musiknoten bzw. Partitur erscheinen. In diesem Fall können andere Submodi, wie Typographie, Layout und Farben, angewendet werden (vgl. Stöckl 2004: 13). Darüber hinaus kann dieser Modus entweder ein zentraler Bestandteil des Geschehens sein und ist dann auch den Protagonist:innen zugänglich (z.B. ein von einer Protagonist:in gesungenes Lied), oder er ist lediglich den Rezipient:innen des audiovisuellen Texts vorbehalten (z.B. eine Szene untermalende Hintergrundmusik). Da die

Bedeutungsinterpretation dieses Hauptmodus sich sowohl durch die Musik selbst als auch durch deren Text vollziehen kann, ist es nicht unbedingt nötig, dass Rezipient:innen den Liedtext verstehen – die Auslegung kann je nach Person variieren und muss auch nicht jener des:der Produzent:in entsprechen (vgl. Pérez-Gonzales 2015: 208f.).

Audiovisuelle Translator:innen haben in Hinblick auf die Vermittlung von Musik in den meisten Fällen mit Liedtexten zu tun. Das Zusammenspiel von musikalischen und verbalen bedeutungsgebenden Elementen lässt hierbei eine Vielzahl an Strategien zu, welche von dem Hauptziel des Textes abhängen. In einer Situation, in der die Bedeutungsübertragung des gesungenen Textes im Vordergrund steht, etwa in Konzertprogrammen, kann eine bloße, mehr oder weniger wörtliche, Übersetzung des Textes bevorzugt werden (vgl. Franzon 2008: 378). In anderen Fällen wird der Struktur des Liedes Vorzug gegeben: „If the tune is the most important component of the music-based ensemble, on the other hand, the lyrics will normally be rewritten to suit the tune’s structural constraints“ (Pérez-Gonzales 2015: 209). Wenn die Singbarkeit oder Vorführbarkeit des Songs genauso wichtig sind wie der Text, kann auch die Melodie geändert werden, um besser mit dem Text zu harmonieren (vgl. Franzon 2008: 381).

Da Musik ein wichtiges Instrument für die Narration ist (vgl. Garwood 2006: 93), ist die Entscheidung, wann Liedtexte übersetzt oder untertitelt werden sollen, ebenfalls von großer Bedeutung. In den meisten Synchronisationsländern bleiben Songs in der Originalsprache bestehen, und auch in Untertitelungsländern gibt es nicht immer ein erkennbares Muster zu Übersetzung oder Nichtübersetzung von Songtexten. Zum Einfluss von nonverbalen filmischen Bedeutungsträgern führte Desilla (2009) eine Studie durch, in der sie untersuchte, ob und inwiefern britische und griechische Rezipient:innen die implizierte Bedeutung in der jeweils englischsprachigen Originalfassung und der auf Griechisch untertitelten Version des Films „Bridget Jones – Schokolade zum Frühstück“ (2001) verstehen können. Genauer wollte sie herausfinden, ob das griechisch-sprachige Publikum die Bedeutung schlechter versteht, wenn keine Untertitelung in der Zielsprache vorhanden ist. Es stellte sich heraus, dass genau das der Fall war – lediglich ein:e Proband:in konnte die versteckte Bedeutung entschlüsseln und erklären, wieso genau dieser Song in diesem Kontext verwendet wurde (vgl. Desilla 2009: 301, zit. n. Pérez-Gonzales 2015: 210ff.). Das Nichtbeachten des Zusammenspiels zwischen Sprache und Musik kann also dazu führen, dass multimodale Texte nur oberflächlich und eventuell sogar unzureichend verstanden werden (vgl. Di Giovanni 2008: 310).

Doch auch ohne zusätzliche Verwendung eines weiteren Modus kann Musik ein wichtiges bedeutungsgebendes Mittel sein. Sie hat oftmals großen Einfluss auf die Rezeption von Charakteren und Settings, wird jedoch selten im audiovisuellen Translationsprozess

verändert und an das Zielpublikum angepasst. Hörende Rezipient:innen, die die interlingual untertitelte oder synchronisierte Fassung eines Films in der Zielsprache ansehen, können den Originalsoundtrack immer noch hören, auch wenn der Liedtext nur wörtlich übersetzt ist, und daher die Stimmung interpretieren, die durch die Musik impliziert wird. Bei einem hörgeschädigten Publikum ist es darum umso wichtiger, die multimodalen Bausteine so genau wie möglich zu beschreiben und ihren Einfluss auf das Gesamtbild richtig einzuschätzen (vgl. Pérez-Gonzales 2015: 211f.).

Ein passendes Beispiel für Musik als bedeutungsgebendes Stilmittel ist der Film „Uhrwerk Orange“ (1971), welcher mit der Dissonanz zwischen Musik und Bild spielt. Regisseur Stanley Kubrick setzt in Szenen, in denen körperliche Gewalt zu sehen ist, verspielte und unbeschwerte Musik ein, um die Hauptfigur Alex zu charakterisieren und die Handlungen beinahe unschuldig wirken zu lassen. Die Musik ist dabei, so Wurm (2007), nicht nur Stilmittel für die Handlung des Films, sondern soll durch die Inkongruenz von visuellen und auditiven Stimuli auch bei den Zuseher:innen Verwirrung auszulösen. Wurm (2007) stellte sich nun die Frage, ob ein hörgeschädigtes Publikum die Komplexität von Alex‘ Charakter genauso verstehen könnte, vor allem, wenn diese sich hauptsächlich durch die Musik darstelle. Die damals geltenden Untertitelungskonventionen, welche vorschlugen, den Liedtext in den Untertiteln einzublenden, würden nicht reichen, um die Stimmung zu reproduzieren. Um das Erlebnis von Filmen auch hörgeschädigten Personen zugänglich zu machen, müsse man die traditionelle Translationswissenschaft verlassen und sich eher Bereichen wie Semiotik und Psychologie zuwenden (vgl. Wurm 2007: 133f.).

Zusammenfassend lässt sich Folgendes sagen: für die audiovisuelle Translation, insbesondere die Untertitelung für Hörgeschädigte und Gehörlose, sind vor allem jene Fälle herausfordernd, in denen die Bedeutungsschaffung des multimodalen Textes sich hauptsächlich durch die Modi Ton, Musik und Sprache vollzieht. Auch auf Inkongruenz zwischen visuellen und auditiven Elementen ist bei der SDH-Untertitelung besonders zu achten, damit gehörlose Rezipient:innen den Inhalt und die Handlung des multimodalen Produktes im gleichen Ausmaß nachvollziehen können wie das hörende Publikum.

1.1.2 Verarbeitung von Multimodalität

Wie bereits zu Beginn dieses Kapitels erwähnt, müssen Rezipient:innen unbewusst intermodale Verknüpfungen herstellen, um das Gesamtbild multimodaler Produkte erkennen zu können. Laut Baldry & Thibault (2006) ist hierfür das *resource integration principle* (Prinzip der

Ressourcenintegration) von zentraler Bedeutung. Multimodale Texte beruhen nicht auf dem Nebeneinanderstellen einzelner Modi, sondern auf der Integration dieser, sodass ein komplexes Ganzes erschaffen werden kann, welches nicht nur als Summe seiner Einzelteile gesehen wird (vgl. Baldry & Thibault 2006: 18). Diese Integration ist jedoch nicht ganz unkompliziert, da jeder Modus eigene bedeutungstragende Eigenschaften besitzt. Stöckl (2004) beschreibt etwa, dass Sprache die sogenannte zweifache Gliederung besitzt, da sie aus bedeutungstragenden und bedeutungsunterscheidenden Bauteilen besteht. Bei Bildern besteht diese Gliederung allerdings nicht – es lässt sich nicht erklären, wieso Pixel mehr Bedeutung tragen, wenn sie kombiniert werden (vgl. Stöckl 2004: 16).

Darüber hinaus erfordert jeder Modus eine bestimmte kognitive Orientation. Sprache als linearer Modus basiert auf dem Zusammenführen von Zeichen zu einem Satz und muss für die weitere Verarbeitung zuerst analysiert werden, während Bilder sofortige Reizeinströmung auslösen und daher auch schneller verarbeitet werden können (vgl. Stöckl 2004: 16f.). Die Integration der verschiedenen bedeutungstragenden Elemente kann also als „multidirectional investment of meaning“ (Yin Yuen 2004: 176) verstanden werden. Da der Modus Bild mehrdeutig sein kann, braucht es für die Kommunikation spezifischer Inhalte oft die Kombination mit anderen Modi (z.B. Sprache). Das bedeutet wiederum, dass der Modus Sprache in einigen Fällen eine dominante Rolle einnehmen kann – ein Umstand, den Yin Yuen (2004) als *contextual propensity* des multimodalen Produkts bezeichnet. Dies wiederum kann den Interpretationsrahmen der Rezipient:innen einschränken (vgl. Yin Yuen 2004: 189). Pérez-Gonzales (2015: 222) beschreibt dies wie folgt: „[T]he more explicit the film director’s communicative intentions are, the more limited is viewers’ capacity to construct their own interpretation.“

1.2 Multimodalität in der Translationswissenschaft

Die Multimodalitätsforschung hat sich in den letzten 20 Jahren aus einer Vielfalt an zusammenhängenden interdisziplinären Forschungsgruppen entwickelt, die festgestellt haben, dass Kommunikation (und somit auch Translation) mehr ist als nur die sprachliche Übertragung von Inhalten. Vor allem die technologischen Entwicklungen der letzten Jahrzehnte führten zu einem gesteigerten (wissenschaftlichen) Interesse an Multimodalität. Kommunikation vollzieht sich über mehr als nur den sprachlichen Modus – so kombinieren etwa Memes geschriebene Sprache mit Bildern, während Vlogs sowohl Bild, Sprache, Text und Musik beinhalten. Der Begriff Multimodalität wird also verwendet, um verschiedenste Ansätze aus den Sozial- und

Kulturwissenschaften sowie aus dem Ingenieurwesen und mittlerweile der Künstlichen Intelligenz zu beschreiben. Dabei vereinen die aufgestellten Theorien, so Jewitt et al. (vgl. 2016: 3), immer folgende drei Aspekte:

1. Bedeutung entsteht durch verschiedene semiotische Ressourcen mit jeweils eigenen Möglichkeiten und Grenzen.
2. Die Bedeutungsgebung geht mit der Produktion multimodaler Gesamtprodukte einher.
3. Zur Untersuchung von Bedeutung müssen alle semiotischen Ressourcen miteinbezogen werden, die Teil eines multimodalen Ganzen sind.

Natürlich gehen all diese verschiedenen Theorien mit unterschiedlichen Traditionen und Konventionen einher; nichtsdestotrotz ist das Feld für viele Bereiche Wissenschaft mittlerweile so relevant, dass einschlägige Journals und Forschungsserien zu Multimodalität geschaffen wurden (vgl. Boria & Tomalin 2020: 11f.).

1.2.1 Der Weg zur Multimodalität

Obwohl seit jeher über das Übersetzen philosophiert und theoretisiert wird – viele westliche Wissenschaftler:innen sehen den Ursprung etwa bei Cicero oder Horaz –, fand dieses Feld erst im 20. Jahrhundert Eingang in die Wissenschaft. Die Entwicklung der Translationswissenschaft wurde von vielen Faktoren vorangetrieben, etwa der Notwendigkeit translatorischer Prozesse bei den Nürnberger Prozessen nach dem Zweiten Weltkrieg, der Entwicklung der Computertechnologie im Kalten Krieg oder auch dem Aufkommen der Sprachwissenschaft als relevantes wissenschaftliches Feld. Die ersten Arbeiten in diesem Feld befassten sich mit der Übertragung geschriebener Sprache und Äquivalenztheorien (z.B. Nida 1964). Bis in die 1960er-Jahre wurde die Translationswissenschaft strikt als Teilbereich der Sprachwissenschaft angesehen; in den 1970ern jedoch etablierte sie sich als eigenes wissenschaftliches Feld und bediente sich nun auch Theorien außerhalb der Linguistik. James Holmes führte schließlich im Jahr 1972 den Begriff *translation studies* ein, und in der Festschrift zur Kreation dieses neuen Feldes hieß es, dass die Translationswissenschaft Übersetzung in all ihren Ausprägungsarten untersuchen wolle. Die Translationswissenschaft sah sich seit jeher als interdisziplinäres Feld, welches nicht nur Akademiker:innen, sondern auch professionelle Übersetzer:innen gleichermaßen inkludierte (vgl. Boria & Tomalin 2020: 7ff.).

Roman Jakobson gilt als einer der Vorreiter der „Eröffnung“ der Translationswissenschaft gegenüber anderen, von den prototypischen Übersetzungstypen abweichenden Bereichen. Jakobson (1959) unterschied drei Arten der Übersetzung –

interlingual, intralingual und intersemiotisch. Die intralinguale Übersetzung (auch *translation proper*) ist hierbei die prototypische Vorstellung von Übersetzung von Inhalten aus einer Sprache in die andere, während die intralinguale Übersetzung (*rewording*) sich innerhalb einer Sprache vollzieht, zum Beispiel, indem ein Text für eine neue Zielgruppe bearbeitet wird. Bei der intersemiotischen Übersetzung, auch *Transmutation* genannt, wird das Zeichensystem gewechselt (vgl. Jakobson 1959: 233). Hiermit ist ein wichtiger Anknüpfungspunkt an die Multimodalität gegeben, da etwa verbale Zeichen durch nicht-verbale Zeichen ausgedrückt werden. Barrierefreie Kommunikation, wie Audiodeskription oder Untertitelung für Hörgeschädigte, zählt in vielen Fällen zu dieser Form der Übersetzung.

Zu den ersten Translationswissenschaftler:innen, die sich mit Multimodalität beschäftigten und diese auch so benannten, zählen Kress & van Leeuwen (2001), welche zeigten, dass prestigeträchtige Formen der Kommunikation in westlichen Kulturen zumeist monomodal waren:

The most highly valued genres of writing (literary novels, academic treatises, official documents and reports, etc.) came entirely without illustration, and had graphically uniform, dense pages of print. Paintings nearly all used the same support (canvas) and the same medium (oils), whatever their style or subject. In concert performances all musicians dressed identically and only conductor and soloists were allowed a modicum of bodily expression. (Kress & van Leeuwen 2001: 1)

Die Monomodalität ist hierbei, so O’Sullivan (2013), jedoch zum Teil eine Sache der Auslegung: Da der Großteil der Translationstheorie – und somit auch das wissenschaftliche Vokabular – sich bis zu diesem Zeitpunkt auf geschriebene Texte fokussierte, wurde auch Translation zumeist als Übersetzung geschriebener Texte gesehen. Die Ressourcen, die bei der Erstellung der Texte verwendet wurden, sowie andere bedeutungsgebende Modi wurden meistens nicht genauer (wissenschaftlich) betrachtet. Diese Auffassung spiegelt sich auch bei Gambier (2006) wider, welcher, wie zu Beginn dieses Kapitels schon beschrieben, behauptet, dass Bedeutung immer multisemiotisch ist (vgl. Gambier 2006: 6).

Obwohl die Wissenschaft im Allgemeinen Multimodalität lange nicht als wichtigen kommunikativen Faktor beachtet hat, war diese in den meisten Kommunikationskontexten der Menschen vorhanden. So sind mittelalterliche Manuskripte oft multimodal angelegt, da sie kalligraphische und illustrative Aspekte beinhalten. (vgl. Jones 2013: 1). Ein Beispiel hierfür ist die „*Practica Chirurgiae*“ aus dem 14. Jahrhundert, deren Illustrationen häufig auch Wortspiele enthielten. So hätte die Eule, die neben dem Begriff für ein Krebsgeschwür dargestellt war – auf Lateinisch tragen beide die Bezeichnung „*bubo*“ – Lai:innen das Verständnis der Fachterminologie erleichtert. Übersetzer:innen der damaligen Zeit hätten diese

Kombination aus Wort und Bild für die Zielsprache entziffern müssen, was eine große Herausforderung darstellen konnte (vgl. O’Sullivan 2013: 3).

Die Multimodalitätstheorie ist aufgrund ihrer soziologischen, funktionalen und medialen Aspekte für viele wissenschaftliche Disziplinen relevant. Da die Translationswissenschaft allerdings viele Jahre lang stark an die Sprachwissenschaft als Mutterdisziplin gebunden war, wurde Multimodalität lange nicht integriert. Darüber hinaus war der Untersuchungsschwerpunkt in den frühen Jahren der Translationswissenschaft ein monomodaler: Die maschinelle Übersetzung der 1950er-Jahre fokussierte sich auf Sprachsysteme und auf die größtmögliche Invarianz von Wort- und Satzeinheiten zwischen zwei Sprachen. Diese Grundeinstellung war auch ein Grund dafür, dass etwa Roman Jakobsons (1959) Theorie, nach der wir auch kommunikative Mittel außerhalb von Sprache verwenden, lange weitgehend ignoriert und erst zu einem späteren Zeitpunkt für die Translationswissenschaft relevant wurde. Mit der Verlagerung hin zu einem kultursensiblen, soziologisch motivierten und medienbewussten Zugang wurde das Konzept immer wichtiger in der Translationswissenschaft; vor allem die wissenschaftliche Untersuchung audiovisueller Translation war ein wichtiger Schritt (vgl. Kaindl 2020: 52ff.).

1.2.2 Zugrundeliegende Theorie: (Multimodale) Sozialesemiotik

Die Multimodale Sozialesemiotik (MSS) war eine der einflussreichsten Theorien für die Entwicklung einer Theorie der multimodalen Kommunikation. Sie basiert auf Michael Hallidays Theorie der Sozialesemiotik (1978). Diese beschreibt, dass Sprache innerhalb eines soziokulturellen Kontextes untersucht werden sollte, in dem die Kultur durch semiotische Begriffe interpretiert wird. Bedeutungsgebung ergibt sich also durch soziale Praktiken in einer Gemeinschaft. Halliday identifiziert hierbei drei verwandte Faktoren – *field*, *tenor* und *mode* –, anhand derer Semiose untersucht werden kann. Field beschreibt hierbei den Inhalt der Diskussion, tenor die soziale Beziehung zwischen den Akteur:innen und mode den Kommunikationskanal (z.B. Text). Mithilfe dieser drei Faktoren kann jegliche Art der Kommunikation charakterisiert werden, und wenn mehr als ein Modus verwendet wird, wird die Kommunikation multimodal (vgl. Boria & Tomalin 2020: 12; Halliday 1978).

Kress & van Leeuwens Werk „Multimodal discourse: the modes and media of contemporary communication“ (2001) war ausschlaggebend für die Diskussion über eine multimodale Kommunikationstheorie in der Translationswissenschaft. Die Autoren definierten Modi als „semiotic resources which allow the simultaneous realisation of discourses and types

of (inter)action“ (Kress & van Leeuwen 2001: 21). Demnach sind Modi nicht in erster Linie Produkte, sondern kulturelle Prozesse, welche sich als Diskurse erkennbar machen und deren Funktionen Texte im Verhältnis mit anderen Modi bilden. Kress & van Leeuwen (vgl. 2001: 20) beschreiben Multimodalität als das Zusammenführen verschiedener semiotischer Modi zu einem Produkt. Die Multimodalitätstheorie erkennt also die Tatsache an, dass verschiedene Modi in Zusammenhang miteinander stehen (vgl. Kaindl 2013: 258).

Die Sozialesemiotik, welche Kress & van Leeuwen (2001) als theoretischen Rahmen für ihre Multimodalitätstheorie annahmen, basiert auf vier grundlegenden Faktoren: den Dimensionen *Diskurs* und *Design* sowie den Phasen *Produktion* und *Distribution*. Während sich die Dimensionen auf die verschiedenen Modi und ihr Design beziehen, werden die Phasen vor allem mit den verwendeten Medien in Verbindung gebracht. Diskurse sind in dieser Theorie sozial situierte Formen von Wissen, welche von Genre, Modus und Design abhängig sind. Kress & van Leeuwen (2001) definieren Design als Mittel, um Diskurse im Kontext einer bestimmten kommunikativen Situation zu realisieren. Im Design wird also die Form des Textes konkretisiert und in der Produktionsphase realisiert, bevor der Text in der Distributionsphase technisch vervielfältigt wird (vgl. Kaindl 2013: 258). Darüber hinaus beleuchtet dieses Konzept die soziokulturelle Dimension: Es geht darin nicht nur um die Dynamik der Bedeutungsschaffung, sondern auch um deren Bindung an sozial etablierte Regeln (vgl. Kaindl 2020: 52).

Es ist hierbei wichtig zu erwähnen, dass die Einbettung von Modi in einen sozial-kommunikativen Kontext bereits seit den 1970er-Jahren in der Semiotik bekannt war, jedoch ermöglicht Kress & van Leeuwens Konzept ein besseres Verständnis der funktionalen Interaktion zwischen semiotischen Ressourcen und deren Darstellungsformen (vgl. Kaindl 2020: 52).

1.2.3 Multimodalitätsforschung heute

In der heutigen Forschung zu Multimodalität in der Translationswissenschaft existieren viele verschiedene Schwerpunkte, welche den interdisziplinären Charakter des Feldes weiter unterstreichen. So beschäftigt sich etwa ein Forschungsstrang mit Layout und Typographie multimodaler Texte, zum Beispiel in Bezug auf die Bibelübersetzung, bei der strategisch romanische und gotische Schriftarten verwendet wurden (vgl. O’Sullivan 2013: 5-6). Diese Aspekte spielen auch im Kontext der audiovisuellen Translation eine große Rolle, etwa bei Fansubbing (vgl. Pérez-González 2007) oder auch bei Untertiteln für Gehörlose (vgl. Neves

2005). Auch multimediale Räume, in denen multimodale Produkte dargeboten werden (z.B. Theater, Museen, Kinos), werden untersucht. Der Fokus liegt etwa auf dem Einfluss von Kameraperspektiven (vgl. Pérez-González 2007) oder Lichtverhältnissen (vgl. Maszerowska 2012) auf die Übersetzung von Filmen und Serien. Im Museumskontext geht es unter anderem um Objekt-Text-Interdependenz, ästhetische Kohärenz oder intertextuelle Beziehungen innerhalb einzelner Ausstellungen (vgl. Sell 2015; Neather 2012; Liao 2018).

Die audiovisuelle Translation nimmt in der Multimodalitätsforschung eine große Rolle ein. Heute beschäftigt sich die Forschung zum Beispiel mit Untertiteln für Gehörlose und Hörgeschädigte, Respeaking und Audiodeskription, welche als Formen der Barrierefreien Kommunikation Personen mit Sinnesbehinderungen die Teilnahme an multimodalen Produkten ermöglichen sollen³. Forschungsschwerpunkte liegen hier zum Beispiel auf der Charakterisierung von handelnden Figuren in Serien und Filmen aufgrund paraverbaler Merkmale und auf deren Übersetzung (vgl. Queen 2004).

In der gegenwärtigen Forschung zu Multimodalität in der Translation ist darüber hinaus die Frage vorherrschend, wo Translation endet und Adaptation beginnt. Einige Forscher:innen betonen, dass, sobald Translation so breit gefasst wird, dass sie die Grenzen eines interlingualen Transfers überschreitet, auch die Grenzen zwischen Adaptation und Translation nicht mehr eindeutig sind. Beide Arten der Textschaffung erfordern Rekonstruktion von Texten in einem System, das sich von jenem unterscheidet, in dem sie ursprünglich erschaffen wurden (vgl. Weissbrod & Kohn 2019: 4). Azenha & Moreira (2012) widmen gar einen ganzen Artikel der Frage: „Don’t translators ‚adapt‘ when they ‚translate‘?“.

Viele Wissenschaftler:innen betonen außerdem, dass Multimodalität derzeit noch keinen allgemein gültigen theoretischen Rahmen und kein einheitliches Set an Analysemöglichkeiten besitzt. Zurzeit werden kognitive und neurolinguistische Translationsmodelle eingesetzt, um zu verstehen, welchen Einfluss die kognitive Salienz einzelner Modi auf Entscheidungen von Translator:innen haben kann. Zum Beispiel beeinflusst der Soundtrack die Aufmerksamkeit der Zuseher:innen beim Verständnisprozesses eines multimodalen Produktes (vgl. Pérez-González 2015). Um solche Fragestellungen zu beantworten, stehen neben der multimodalen Transkription (vgl. Taylor 2004; Taylor 2016) und multimodalen Korpora (vgl. Soffritti 2018) auch Eye-Tracking-Studien zur Verfügung, welche als sehr effektiv gelten: Das Sehverhalten der Teilnehmer:innen gibt Aufschluss

³ Diese Themen werden in Kapitel 2.2 und Kapitel 3 genauer beleuchtet.

darüber, wie sie einen multimodalen Text verarbeiten und verstehen, und kann so auch für Translator:innen hilfreich sein (vgl. Kruger 2012).

1.2.4 Die Analyse multimodaler Texte

Sowohl Multimodalität als auch multimodale Kohäsion audiovisueller Texte wurden bereits in vielerlei Hinsicht (translations-)wissenschaftlich diskutiert – dennoch gibt es bisher keinen einheitlichen Rahmen für deren Analyse. In vielen Fällen werden Modelle aus anderen Wissenschaften verwendet, zum Beispiel aus der Filmtheorie bei audiovisuellen Produkten oder Musiktheorien bei der Übersetzung von Musicals und Songs. Viele Expert:innen der Translationswissenschaft sehen diese Tatsache als kritisch an, da, wie etwa Kaindl (2013) oder Gambier (2006) betonten, diese Konzepte nicht das hauptsächliche Untersuchungsinteresse, nämlich den Transfer über Kulturbarrrieren hinweg, miteinbeziehen. Aus diesem Grund brauche es Untersuchungsmethoden, die an die Anforderungen der Translationswissenschaft angepasst sind (vgl. Kaindl 2013: 264).

Seit dem Eingang von multimodalen und audiovisuellen Produkten in die Translationswissenschaft gab es diesbezüglich unterschiedliche Auffassungen und Theorien, welche in weiterer Folge auch Auswirkungen auf die anzuwendenden Untersuchungsmethoden hatten. Als die Übersetzungswissenschaft sich zum Beispiel vor allem auf jene Aspekte fokussierte, die audiovisuelle Texte von anderen Texttypen unterschied, basierte auch die Untersuchungsmethodik auf den Einschränkungen für Translator:innen beim Übersetzen audiovisueller Produkte (vgl. Titford 1982; Mayoral et al. 1988).

Remael & Reviers (2018) identifizieren drei Analysemodelle für multimodale Texte, welche sich im Laufe der Zeit etabliert haben und etwa im Kontext der Filmübersetzung Verwendung finden. Das erste Modell, das hier näher beschrieben werden soll, ist van Leeuwens (2005) Konzept aus der Sozialemiotik. Dieses Modell basiert auf dem Zusammenspiel von vier Faktoren, welche die Integration und das gleichzeitige Auftreten verschiedener semiotischer Ressourcen steuern: Rhythmus, Komposition, Informationsverknüpfung und Dialog. Während der Rhythmus den Geschehnissen Kohärenz und Struktur gibt, erfüllt die Komposition dieselbe Funktion für den Raum. Die Informationsverknüpfung wiederum bezieht sich auf kognitive Zusammenhänge einzelner Informations-Items oder auf Kausalverbindungen zwischen Sprache und Bild. Der Faktor Dialog beinhaltet die Strukturen dialogischer oder musikalischer Äußerungen in multimodalen Texten (vgl. van Leeuwen 2005). Van Leeuwens Modell kann jedoch sehr breit ausgelegt

werden und ist als solches nicht als systematischer Analyserahmen für multimodale Texte ausgelegt, weshalb viele Wissenschaftler:innen auf seinem Modell aufbauend eigene, spezifischere Analysemodelle geschaffen haben (vgl. Remael & Reviers 2018: 260f.).

Eines davon ist Royces (2007) Modell, welches die Verknüpfung der einzelnen Faktoren mit Begriffen aus der Sprachwissenschaft, genauer aus dem Konzept der linguistischen Kohäsion, erweitert. Die Verknüpfung vollziehe sich demnach als Zusammenspiel der Faktoren Referenz, Wiederholung, Synonymie, Kollokation und Ganzes-Teil-Beziehungen. Auf diesem Weg kann auch die Verbindung nicht-verbaler Modi beschrieben werden (vgl. Royce 2007: 103). Remael & Reviers (2015) sehen diese Faktoren jedoch als nicht ausreichend an, da die Informationsvermittlung sich auch auf implizitem Weg oder durch dialogische Interaktion, ähnlich wie bei van Leeuwen (2005), vollziehen kann, und schlagen daher den zusätzlichen Faktor Komplementarität vor. Dieser beschreibt die Beziehung zwischen zwei Zeichen, die „einfach nur“ zur gleichen Zeit auftreten. Auf diese Weise kann zwischen lexikalischer Kohäsion und Komplementarität unterschieden werden, auch wenn, so Remael & Reviers (vgl. 2015: 56), dies einen gewissen Interpretationsspielraum zulasse.

Das letzte Analysemodell, das hier beschrieben werden soll, basiert auf dem Konzept der Sozialesemiotik und soll die zuvor genannten Modelle zu einem systematischen Analysemodell vereinen. Tseng (2013) ist der Auffassung, dass Rezipient:innen die Narration eines Filmes aufgrund von Bedeutungsmustern erkennen, welche sie durch vorherige Erfahrungen (etwa das Ansehen von Filmen oder sonstige Lebenserfahrungen) erworben haben. Dieses Verständnis spiegelt sich vor allem in Narrationstheorien wider, wie sie auch in Kapitel 3.3.6 genauer beschrieben werden. Rezipient:innen eines Filmes bedienen sich laut Tseng (2013) vier Faktoren, um eine kohärente Filmnarration zu bilden: Charaktere, Objekte, Settings und Aktionen. Sie bilden also die Identitäten von Charakteren und Objekten, indem sie kohäsive Verkettungen erschaffen, welche ihre Interpretation der Narration anleiten. Dieser Ansatz ist laut Remael & Reviers (2018) besonders gut für die Untersuchung audiovisueller Texte geeignet, da die Kohäsion sich in diesen Texten sowohl durch explizite als auch implizite Bedeutungsbeziehungen vollzieht (vgl. Remael & Reviers 2018: 261f.).

Multimodale Kohäsion ist eines der Hauptaugenmerke in der audiovisuellen Translation und Media Accessibility. Da in diesem Bereich der Übersetzung ein Modus verändert wird, kann das auch Einfluss auf die Interaktion mit anderen Modi haben. Übersetzer:innen audiovisueller Produkte müssen daher besonders darauf achten, die multimodale Kohäsion aufrechtzuerhalten und ein kohärentes und kohäsives Endprodukt zu schaffen. Bei barrierefreier

Kommunikation wie Audiodeskription oder SDH ist darauf besonderes Augenmerk zu legen (vgl. Remael & Reviers 2018: 262).

2. Barrierefreiheit und Translationswissenschaft

2.1 Grundlagen der Barrierefreien Kommunikation

2.1.1 Definition von Barrierefreiheit und Barrierefreier Kommunikation

Barrierefreiheit (*Accessibility*) wird laut dem Bundesgesetz über die Gleichstellung von Menschen mit Behinderungen (Bundes-Behindertengleichstellungsgesetz, BGStG) folgend definiert:

Barrierefrei sind bauliche und sonstige Anlagen, Verkehrsmittel, technische Gebrauchsgegenstände, Systeme der Informationsverarbeitung sowie andere gestaltete Lebensbereiche, wenn sie für Menschen mit Behinderungen in der allgemein üblichen Weise, ohne besondere Erschwernis und grundsätzlich ohne fremde Hilfe zugänglich und nutzbar sind. (§6 Abs 5 BGStG).

Eine Behinderung liegt laut BGStG dann vor, wenn eine nicht nur vorübergehende, also länger als sechs Monate andauernde, körperliche, geistige oder psychische Funktionsbeeinträchtigung oder Beeinträchtigung der Sinnesfunktionen geeignet ist, die Teilhabe am Leben in der Gesellschaft zu erschweren (vgl. §3 BGStG).

Ein Aspekt, den auch Maaß & Rink (2018) im deutschen Behindertengleichstellungsgesetz kritisieren, ist die Formulierung „ohne besondere Erschwernis“, welcher sich auf verschiedene Dimensionen beziehen könne: etwa die körperliche oder geistige Anstrengung, aber auch auf den finanziellen Aufwand. Diese Vagheit erschwere das Geltendmachen von Ansprüchen, weshalb auch das tatsächliche Erlangen barrierefreier Angebote für Menschen mit Behinderung in der Praxis oft nicht einfach ist (vgl. Maaß & Rink 2018: 18f.).

Die UN-Behindertenrechtskonvention, die in Österreich seit 26. Oktober 2008 in Kraft ist, empfinden Maaß & Rink (2018) als aussagekräftiger in Bezug auf den Aspekt der Kommunikation. Sie besagt:

Kommunikation [schließt] Sprachen, Textdarstellung, Brailleschrift, taktile Kommunikation, Großdruck, barrierefreies Multimedia sowie schriftliche, auditive, in einfache Sprache übersetzte, durch Vorleser zugänglich gemachte sowie ergänzende und alternative Formen, Mittel und Formate der Kommunikation, einschließlich barrierefreier Informations- und Kommunikationstechnologie, ein. (UN-BRK 2008 Art. 2)

Der Begriff *Sprache* umfasst hierbei sowohl gesprochene als auch nicht-gesprochene und Gebärdensprachen (vgl. UN-BRK 2008 Art. 3).

Da die UN-Behindertenrechtskonvention jedoch, wie Maaß & Rink (vgl. 2018: 19f.) für Deutschland feststellen, auch in Österreich nicht mit der tatsächlichen Rechtsgrundlage korreliert, sei sie vielmehr ein Ausgangspunkt für Forschung und Praxis.

Basierend auf diesen Gesetzen bzw. Konventionen schlagen Maaß & Rink (2018: 20) folgende Definition für den Begriff *Barrierefreie Kommunikation* vor: „Barrierefreie Kommunikation umfasst alle Maßnahmen zur Eindämmung von Kommunikationsbarrieren in unterschiedlichen situationalen Handlungsfeldern.“ Kommunikationsbarrieren können in diesem Fall in Hinblick auf Sinnesorgane und/oder kognitive Voraussetzungen der Kommunikationsteilnehmer:innen bestehen, doch auch die Anforderungen – seien diese sprachlich, fachsprachlich, fachlich, kulturell oder medial – der Texte an die Rezipient:innen. Diese Barrieren entstehen immer in Situationen, in denen das Kommunikationsangebot nicht in der erforderlichen Weise an die Zielsituation und das Zielpublikum angepasst ist (vgl. Maaß & Rink 2018: 20).

Natürlich kann auch Barrierefreie Kommunikation so weit bzw. eng wie möglich bzw. nötig gefasst werden. Während einige Autor:innen auch die Kommunikation zwischen Expert:innen und Lai:innen ohne kommunikative Beeinträchtigungen durch Barrieren bedroht sehen, sind andere der Auffassung, dass sie sich nur auf eine Rezipient:innenschaft bezieht, welche etwa durch Behinderung oder divergierende Bildungschancen angepasste Texte benötigt (vgl. Maaß & Rink 2018: 21).

Neben all diesen definitorischen Überlegungen zum Thema Barrierefreiheit darf jedoch nicht vergessen werden, dass vollständige Barrierefreiheit im wörtlichen Sinne, so viele Expert:innen, eine Utopie ist und auch bleiben wird. Die wirkliche Barrierefreiheit für alle Nutzer:innengruppen wird nicht bzw. kaum realisierbar sein, da allein die Anforderungen für all diese Gruppen sich widersprechen können. Die Ressourcen für die Erstellung Barrierefreier Kommunikation sind außerdem begrenzt, und das Feld ist noch nicht professionalisiert. Es wäre daher generell ratsam, anstelle von „barrierefrei“ eher die Begriffe „barrierearm“ oder auch „barriereärmer“ zu verwenden (vgl. Maaß & Rink 2018: 21). Da aber auch in den Rechtstexten und in der einschlägigen Literatur der Terminus „Barrierefreiheit“ verwendet wird, wird dieser auch in der vorliegenden Arbeit eingesetzt.

2.1.2 Teilhabe und Ableismus

Die eigenständige Teilhabe am öffentlichen Leben ist eines der zentralen Konzepte der Barrierefreiheit und folglich auch der Barrierefreien Kommunikation. Die Teilhabe bzw.

Partizipation ist ein Konzept, welches sich auf Interaktion im gesellschaftlichen Kontext bezieht und auf ein Miteinander von Menschen mit und ohne Behinderung abzielt. Für diesen Zweck braucht es, so Maaß & Rink (2018), barrierefreie Kommunikate, denn nur auf diese Weise haben alle Personen den geforderten eigenständigen Zugang zu Informationen, was wiederum die Voraussetzung für Teilhabe ist. Die Zugänglichkeit von Kommunikaten wird in fünf aufeinanderfolgende Schritte unterteilt: Auffindbarkeit, Wahrnehmbarkeit, Verständlichkeit, Verknüpfungsfähigkeit und Handlungsorientiertheit. Diese Faktoren gelten sowohl aus der Perspektive der Nutzer:innen als auch aus der Perspektive der Kommunikate (vgl. Maaß & Rink 2018: 24).

Schuppener & Bock (2018: 223) betonen, dass „die Frage nach Teilhabe oder Partizipation von Menschen in marginalisierten Lebenslagen als bedeutsam und herausfordernd gleichermaßen [erscheint].“ Partizipation wird häufig synonym mit Teilnahme, Teilhabe, Beteiligung oder Mitbestimmung verwendet und wird oft, so die Autorinnen, unhinterfragt positiv und moralisch überhöht dargestellt. Sie wird in diesem Zusammenhang nicht nur als aktive Beteiligung, sondern gleichzeitig als ein Recht auf Mitsprache, Mitgestaltungsmöglichkeiten und Mitbestimmung definiert (vgl. Schuppener & Bock 2018: 224).

Das Verständnis von Teilhabe ist abhängig von dem in der Gesellschaft vorherrschenden Modell von Behinderung. Jahrzehntlang war in unserer Gesellschaft der Begriff „Behinderung“ gleichgesetzt mit Ohnmacht, Krankheit und Hilflosigkeit, was behinderte Menschen in weiterer Folge als defizitär darstellt. Die Betroffenen erleben also in ihrem Alltag oftmals Misstrauen; Sämtliche Fähigkeiten werden ihnen abgesprochen. Dieses auch heute noch verbreitete, defizitäre Konzept von Behinderung ist ein Ausdruck des *Ableismus*, der bei einem „nichtbehinderten Ideal“ seinen Ausgang nimmt und sich in Praktiken, Ideen und Überzeugungen niederschlägt, welche den behinderten Menschen diskriminieren (vgl. Schuppener & Bock 2018: 222).

Dieses Verständnis von Behinderung kann folglich nicht zu positiven Zielen wie Teilhabe, Gleichstellung oder Selbstbestimmung führen. Es ist daher von großer Bedeutung, dass das Konzept „Behinderung“ von einem ganzheitlichen Modell ausgeht, welches neben den medizinischen Aspekten auch Umwelteinflüsse in Betracht zieht (vgl. Winter 2014: 17).

Bei der Realisierung von Partizipation muss auch berücksichtigt werden, dass der Diskurs um dieselbe immer auch mit dem Anspruch an Demokratie, Gerechtigkeit und Institutionalisierung verbunden ist. Vor allem letzteres Konzept ist für Menschen mit Behinderung zentral, da ein großer Anteil des Lebensalltags von Menschen aller Altersklassen

mit dem Label „behindert“ in Institutionen stattfindet, die wiederum „Orte der Produktion und Reproduktion von Hierarchien und Machtstrukturen“ (Schuppener & Bock 2018: 224) sind. In weiterer Folge könnte, so einige Expert:innen, wahre Partizipation nur dann stattfinden, wenn Machtverhältnisse reflektiert und verändert würden (vgl. Schuppener & Bock 2018: 224f.).

Schuppener & Bock (2018) schlagen daher ein Partizipationskonzept vor, das einerseits die Beteiligung und die Identifikation mit Institutionen, Werten und Kräften der Gesellschaft ist und sich andererseits engagiert und aktiv an demokratischen Strukturen und Prozessen beteiligt. Dieses Konzept „verträgt sich nicht mit der traditionellen Behindertenarbeit, die von einem defekt- und defizitorientierten Denken und Handeln geprägt war und da und dort noch ist, welches Menschen mit Behinderungen als versorgungs-, behandlungs- und anweisungsbedürftige Mängelwesen betrachtet“ (Schwalb & Theunissen 2018: 9). Es müsse daher eine Scheinpartizipation vermieden werden, in welcher Menschen mit Behinderungen nur innerhalb vorgegebener Konzepte und Institutionen handeln können. Außerhalb dieser Strukturen ist ein Agieren für sie nicht möglich (vgl. Schuppener & Bock 2018: 226).

Das Ziel des Paradigmenwechsels in Hinblick auf Behinderung solle also sein, von einer Auffassung als „verminderter Zustand des Menschseins“ (Campbell 2008, zit.n. Schuppener & Bock 2018: 222) und Bevormundung hin zu wahrer Selbstbestimmung zu führen (vgl. Winter 2014: 20).

2.1.3 Kommunikationsbarrieren

Kommunikationsbarrieren können, wie im ersten Teilkapitel bereits skizziert wurde, auf verschiedene Arten existieren. Abhängig von ihrer Ausprägung können diese die Rezipient:innen „den Zugang zum Textgegenstand erschweren, da sie negativ mit der Perzeptibilität und der Verständlichkeit korrelieren“ (Rink 2018: 29). Schubert (2016) und Rink (2018) zählen sieben wahrnehmungs- oder verständnisbasierte Kommunikationsbarrieren auf, die im Folgenden dargestellt werden sollen:

- Eine *Sinnesbarriere* besteht dann, wenn ein zur Aufnahme von Information notwendiger Sinneskanal in seiner Funktion erheblich beeinträchtigt ist, sodass die Rezeption über diesen Kanal nicht möglich ist. Texte können auditiv, visuell oder haptisch rezipiert werden, wobei auch eine Überschneidung dieser Modi möglich ist. Wenn ein erforderlicher Wahrnehmungskanal nicht zur Verfügung steht, welcher zur Aufnahme eines Kommunikats benötigt wird, liegt eine Sinnesbarriere vor. Im Fall der SDH-Untertitelung soll also eine solche Sinnesbarriere behoben werden, indem der

auditive Kommunikationsinhalt über die Untertitelung zugänglich gemacht wird (vgl. Rink 2018: 30).

- Eine *Fachbarriere* entsteht durch Fachlichkeit, welche sich auch in der Sprache bemerkbar machen kann (s. nächster Abschnitt). Der Textgegenstand wird also nicht verstanden, weil „den betroffenen Menschen fachliches, also inhaltliches, Wissen fehlt“ (Schubert 2016: 18). Im Gegensatz zur Sinnesbarriere kann die betroffene Person die Mitteilung also wahrnehmen, aber nicht verstehen (vgl. Schubert 2016: 18).
- Ist eine *Fachsprachenbarriere* vorhanden, so verstehen die Rezipient:innen zwar die Sprache, jedoch nicht die Fachsprache des betroffenen Kommunikats. Es ist dabei möglich, dass Fach- und Fachsprachenbarriere gemeinsam auftreten – sie müssen aber nicht gekoppelt sein. Laut Rink (2018) stellen Fachtexte derzeit eine der größten Herausforderungen der Barrierefreien Kommunikation dar, da sie spezielle Maßnahmen erfordern (vgl. Rink 2018: 30f.).
- Eine *Kulturbarriere* hindert bestimmte Personen(gruppen) daran, ein Kommunikat vollständig zu verstehen, da ihnen das notwendige kulturelle Wissen fehlt. Unter Kultur wird nach Vermeer (1990: 36) „die Menge aller Konventionen einer Gesamtgesellschaft“ verstanden. Es kann sich hierbei um Sach- oder Fachtexte handeln, da diese sich je nach Kultur, zum Beispiel in Hinblick auf Textsortenkonventionen, unterscheiden können (vgl. Schubert 2016: 18).
- Eine *Kognitionsbarriere* liegt vor, wenn „die sprachliche oder inhaltliche Komplexität der Mitteilung“ (Schubert 2016: 18f.) die Rezipient:innen kognitiv überfordert. Auch hier ist das Verständnis der Grund für die Kommunikationsbarriere, da „die Argumentationsstruktur [...] die Verarbeitungskapazität der Rezipient:innen [überfordert]“ (Rink 2018: 31).
- Im Fall einer *Sprachbarriere* kann ein Kommunikat nicht verstanden werden, da es in einer Sprache formuliert ist, die die Rezipient:innen nicht beherrschen (vgl. Schubert 2016: 19). Diese liegt im deutschen Sprachraum vor allem dann vor, wenn Deutsch nicht Erstsprache ist (vgl. Rink 2018: 31).
- Außerdem kann eine *Medienbarriere* bestehen, welche in den Bereichen Codalität, Modalität und Medium auftreten kann. Wenn die Codalität betroffen ist, können sprachliche Informationen, die in bestimmter Weise codiert wurden, nicht aufgelöst werden. Ist die Ebene des Modes eingeschränkt, so ist das zur Informationsaufnahme erforderliche Sinnesorgan nicht funktionstüchtig. Wie bereits bei der Sinnesbarriere

beschrieben, können die Kommunikate auditiv, visuell oder haptisch oder auch in Kombination dieser Modi rezipiert werden, also zum Beispiel audiovisuell. Auch das Trägermedium als Übermittler des Kommunikats kann zur Barriere werden. Dies ist vor allem dann relevant, wenn die Heterogenität der Zielgruppe Barrierefreier Kommunikation in Betracht gezogen wird, „da ihre speziellen Bedarfe und Nutzungsformen den präferierten Zugang auf Inhalte über bestimmte Medien determinieren“ (Rink 2018: 32) (vgl. Rink 2018: 32).

Abhängig von ihrer Gestaltung können diese Barrieren in Texten in unterschiedlich starkem Ausmaß vorhanden sein. Das Anforderungsprofil der Adressat:innenschaft legt fest, ob dieses Ausmaß sich negativ auf Perzeptibilität und Verständlichkeit auswirkt (vgl. Rink 2018: 32).

2.1.4 Arten von Kommunikationseinschränkungen

Die Rezipient:innenschaft von Barrierefreier Kommunikation lässt sich in zwei Gruppen aufteilen, nämlich in jene mit Perzeptionseinschränkungen und jene mit Verständniseinschränkungen.

2.1.4.1 Perzeptionseinschränkungen

Unter Perzeptionseinschränkungen versteht man die sogenannten Sinnesbehinderungen, also Seh- und Hörschädigung. Eine Sehbehinderung wird durch drei Faktoren bestimmt: die Sehschärfe (Visus), die Qualität des Sehvermögens und das Gesichtsfeld, also den wahrgenommenen Bereich ohne Bewegung der Augen. Vor allem Visus und Gesichtsfeld bestimmen den Grad der Sehbehinderung. Laut dem Berufsverband der Augenärzte in Österreich und Deutschland gelten in Österreich folgende Schweregrade der Sehbeeinträchtigung, die sich am verbliebenen Ausmaß der Sehschärfe des besseren Auges orientieren: Als sehbehindert gelten Personen, die trotz optischer Hilfen auf dem besseren Auge über einen Sehrest von nicht mehr als 30 Prozent verfügen, bei hochgradiger Sehbehinderung beträgt dieser Wert weniger als fünf Prozent. Als blind gelten Personen, die über einen Sehrest von weniger als zwei Prozent oder über eine Gesichtsfeldeinschränkung von weniger als fünf Grad verfügen (vgl. Besthelp o.J.). Laut Blinden- und Sehbehindertenverband Österreich wird angenommen, dass etwa 300.000 Menschen in Österreich von Blindheit und Sehbehinderung betroffen sind (vgl. BSVÖ o.J.).

Die zweite Form von Sinnesbehinderungen sind Hörschädigungen. Da in Kapitel 3, welches sich mit SDH-Untertitelung auseinandersetzt, auch die Heterogenität der Zielgruppe

beschrieben wird (Kapitel 3.2.3), soll hier nur eine kurze Skizzierung erfolgen. Die Höreinschränkung wird in Österreich in fünf Gruppen eingeteilt:

- Leichtgradig schwerhörig (Hörverlust von 20 bis 40 dB),
- Mittelgradig schwerhörig (Hörverlust von 40 bis 60 dB),
- Hochgradig schwerhörig (Hörverlust von 60 bis 80 dB),
- An Taubheit grenzend schwerhörig (Hörverlust von 80 bis 90 dB)
- und gehörlos (Hörverlust von mehr als 90 dB).

Der Grad einer Hörminderung wird, ähnlich wie bei der Sicht, an der Hörschwelle des besser hörenden Ohrs bestimmt, wobei auch eine einseitige Hörminderung große Schwierigkeiten bereiten kann. Hinzu kommt, dass Hören nicht mit Verstehen gleichzusetzen ist. Durch den Ausfall bestimmter Hörbereiche können die betroffenen Personen teilweise noch etwas hören, verstehen das Gesagte aber nicht, was eine große körperliche und psychische Belastung bedeuten kann (vgl. Gesundheit.gv.at 2021).

Rink (2018) stellt mithilfe eines Barriereindex, der den Zusammenhang zwischen Barrieretypen und Adressat:innenprofil darstellt, fest, dass „das Anforderungsprofil einer Adressat:innengruppe umso höher ist, je mehr Kommunikationsbarrieren für eine Adressat:innenschaft bestehen“ (Rink 2018: 57). Die Gruppe der (prälingual) Hörgeschädigten weist dabei den höchsten Barriereindex auf, da sie von sämtlichen Barrieren betroffen ist; Diese Gruppe hat besonders ausgeprägte Anforderungen auf perzipierbare und verständliche Texte. Da keine kognitive Einschränkung vorhanden ist, sind die Strategien zum Barriereabbau vor allem sprachlicher und medialer Art.

Außerdem kann die sogenannte Taubblindheit auftreten – eine Behinderung eigener Art, die nicht nur die Addition von Taubheit und Blindheit ist, sondern von einer Schädigung des Hörens und Sehens ausgeht. Da beide Sinne geschädigt sind, können die Ausfälle des einen Sinnes nicht durch den anderen ausgeglichen werden (vgl. ÖHTB 2022). In Österreich sollen laut Forum für Usher-Syndrom, Hörsehbeeinträchtigung und Taubblindheit (USH+TB) rund 200.000 Menschen von Taubblindheit oder anderen Hörsehbeeinträchtigungen betroffen sein (vgl. USH+TB o.J.). Beim Usher Syndrom, auch Morbus Usher genannt, welches der Auslöser für mehr als 50 Prozent dieser Beeinträchtigungen ist (vgl. USH+TB o.J.), tritt zusätzlich zu einer ab Kindesalter bestehenden Schwerhörigkeit im Erwachsenenalter eine Netzhautdegeneration auf. Diese wird durch die sogenannte Retinopathia Pigmentosa, RP, verursacht und führt zum Absterben der Netzhautzellen (vgl. ÖHTB 2022).

Taubblinde Menschen können über den taktil-kinästhetischen Kanal kommunizieren, etwa durch das Lorm-Alphabet, auch Lormen genannt. Hierbei wird durch das Abtasten des Fingeralphabets auf der Hand oder an bestimmten Handarealen der Austausch von Informationen ermöglicht. Rink (2018) stellt für diesen Bereich einen enormen Forschungsbedarf fest, der vor allem die kommunikativen Anforderungen von Taubblinden betrifft (vgl. Rink 2018: 35).

2.1.4.2 Verständniseinschränkungen

Da Personen mit Verständniseinschränkungen nicht zur primären Zielgruppe von SDH-Untertitelungen gehören, sie im Zuge der Barrierefreien Kommunikation dennoch für sich nutzen können, wird diese Zielgruppe hier nur kurz beschrieben.

Zum einen zählen Menschen mit geistiger Behinderung zu dieser Zielgruppe. Hier ist aber zu erwähnen, dass die Bezeichnung „geistige Behinderung“ im Moment Grundlage vieler Diskussionen ist und von der Gruppe selbst als diskriminierend empfunden wird. Aus diesem Grund wählen die betroffenen Menschen oft die Selbstbezeichnung „Menschen mit Lernschwierigkeiten“, um einer Stigmatisierung entgegenzuwirken. Dieser Begriff wird aber in der einschlägigen Literatur für eine andere Personengruppe verwendet, welche im Folgenden auch noch beschrieben wird. Die Vermischung der Konzepte von Lernschwierigkeit und geistiger Behinderung sind, so Rink (2018), jedoch nicht unproblematisch, da juristisch gesehen nur Personen mit „langfristige[n] körperlichen, seelischen, geistigen oder Sinnesbeeinträchtigungen [...], welche sie in Wechselwirkung mit einstellungs- und umweltbedingten Barrieren an der gleichberechtigten Teilhabe an der Gesellschaft hindern können“ (§3 BGG), Anspruch auf Barrierefreiheit und Verständlichkeit haben. Folglich haben Personen mit Lernschwierigkeiten nur dann einen Rechtsanspruch auf zugängliche Kommunikationsangebote, wenn die Dauer der Beeinträchtigung nachgewiesen ist. Die Verwendung des einen oder anderen bestimmten Terminus hat also Einfluss auf den gesetzlich festgelegten Anspruch (vgl. Rink 2018: 38f.).

Der Begriff Lernschwierigkeiten umfasst verschiedene Ausprägungen von Lernproblemen, welche in allgemeine und gravierende Lernschwierigkeiten unterteilt werden können. Die gravierenden Lernschwierigkeiten unterscheiden sich dadurch von den allgemeinen Lernschwierigkeiten, dass es Unterstützung, Förderung und Begleitung zum Überwinden dieser Lernschwierigkeiten bedarf. Gerade Texte in Leichter Sprache könnten den Zugang zur Informationsgesellschaft erleichtern (vgl. Rink 2018: 40).

Auch für Personen mit Arten von Demenz, also degenerativen neurologischen Erkrankungen, welche ein Zellsterben hervorrufen und vor allem im Alter auftreten, ist Barrierefreie Kommunikation eine mögliche Zugangsweise zu Informationen. Dasselbe gilt für Menschen mit Aphasie, einer Erkrankung des zentralen Nervensystems, welche das Sprachzentrum des Gehirns schädigt und das Sprechen, Lesen und Schreiben beeinträchtigen kann (vgl. Rink 2018: 40f.).

Zur Gruppe der Personen mit Verständniseinschränkungen können außerdem Personen mit Deutsch als Zweitsprache gehören, da sie die deutsche Sprache zum großen Teil oder ausschließlich ohne unterrichtliche Unterstützung erworben haben. Gerade zu Beginn eines Aufenthalts können verständnisoptimierte Texte dabei helfen, Informationen zu vermitteln und so einen Anreiz schaffen, die eigenen Kompetenzen auszubauen. Juristisch gesehen hat diese Gruppe keinen Anspruch auf Barrierefreie Kommunikation; gleichzeitig muss sie vor allem am Anfang ihres Aufenthalts oftmals Formulare und Anträge ausfüllen, die schwer verständlich sind und die Rezipient:innen schnell an ihre Grenzen bringen (vgl. Rink 2018: 43).

Zuletzt sind auch Menschen mit Mehrfachbehinderung, zu denen auch die Gruppe der Taubblinden gehört, die bereits beschrieben wurde, ein Teil der Zielgruppe Barrierefreier Kommunikation. Dabei kann jegliche Form der zuvor beschriebenen Arten von Beeinträchtigungen kombiniert mit einer anderen auftreten, also zum Beispiel Demenz und Hörschädigung oder Deutsch als Zweitsprache und geistige Behinderung. Die Ausprägung kann sich außerdem individuell stark unterscheiden. Die Betroffenen sind nicht nur in Wahrnehmung und Kommunikation beeinträchtigt, sondern in ihren allgemeinen Erlebens- und Ausdrucksmöglichkeiten (vgl. Rink 2018: 44f.).

2.1.5 Strategien zum Barriereabbau

Zur Behebung der in Kapitel 2.1.3 vorgestellten Kommunikationsbarrieren gibt es drei Arten von Strategien: die sprachlichen, die medialen und die konzeptuellen.

Als *sprachliche Strategien* werden jene Maßnahmen bezeichnet, die die Verständlichkeit von Kommunikaten auf Wort-, Satz- und Textebene verbessern sollen, etwa die Verwendung von kürzeren, weniger komplizierten Wörtern und Sätzen, die Veranschaulichung von Sachverhalten mit Beispielen oder die Erläuterung von komplexen Konzepten (vgl. Rink 2018: 60).

Mediale Strategien umfassen vor allem die Verbesserung der Perzeptibilität, also die Wahrnehmbarkeit des Textes. Zu diesem Zweck können verschiedene Mittel der

typografischen Gliederung und Vernetzung eingesetzt werden, darunter Hervorhebungen oder Bebilderungen. Für den Fall, dass bestimmte Sinnesorgane als Wahrnehmungskanäle zur Verfügung stehen, müssen alternative Kanäle bedient werden. Dies kann zum Beispiel durch Untertitel, Audiodeskriptionen, Alternativtexte oder Brailleschrift erfolgen (vgl. Rink 2018: 60). Mediale Strategien „zielen damit auf die multimodale Aufbereitung von Text“ (Rink 2018: 61).

Konzeptuelle Strategien wiederum zielen auf die Absenkung der kognitiven Komplexität von Inhalten. Die Reduktion der Komplexität kann einerseits durch sprachliche und mediale Strategien erfolgen, oder die Struktur des jeweiligen Informationsangebots wird verändert. Ein Beispiel hierfür sind Metatexte, welche Audioeinführungen oder Textversionen in Leichter Sprache zugänglich machen. Außerdem können Inhalte das Medium wechseln – so könnte ein Audioangebot etwa grafisch realisiert werden, wie es auch bei SDH-Untertitelungen der Fall ist (vgl. Rink 2018: 61).

Dieser Unterteilung folgend ist die Bereitstellung von SDH-Untertiteln sowohl eine mediale als auch eine konzeptuelle Strategie der Barrierefreien Kommunikation.

2.2 Das Verhältnis von Barrierefreiheit und Translation

Durch den technologischen Fortschritt der letzten Jahrzehnte ist die Welt, wie wir sie heute kennen, zu einer Informationsgesellschaft geworden. Wir beziehen den Großteil der uns verfügbaren Informations- und Unterhaltungsangebote über multimodale Zugänge, etwa über das Internet, Fernsehen, Computer, aber auch in Theatern, Museen oder an anderen öffentlichen Orten. Die Kombination aus Informations- und Kommunikationstechnologie sowie audiovisueller Translation hat es uns ermöglicht, uns jederzeit auf der ganzen Welt zu vernetzen. Die große Problematik an der audiovisuellen Medienwelt besteht allerdings darin, dass sie größtenteils nur „physically able people“ (Díaz Cintas et al. 2007: 11) zugänglich ist. Der Zugang zur Informationsgesellschaft ist aber notwendig, um an den Vorteilen der Globalisierung sowie an wirtschaftlichem und persönlichem Wachstum teilzuhaben. Die Ausgrenzung von diesem Informationsangebot kann verschiedene Gründe haben (Alter, körperliche Probleme, geographische Abgeschiedenheit, Geldmangel) und kann in weiterer Folge zu sozialer Marginalisierung führen, weshalb der Zugang zu Informationen im Jahr 2003 von der Europäischen Union als Menschenrecht festgelegt wurde (vgl. Díaz Cintas et al. 2007: 11f.).

Viele Jahre lang wurde der Terminus „Barrierefreiheit“ beinahe ausschließlich für Hilfestellungen für körperlich behinderte oder eingeschränkte Personen verwendet. Das führte dazu, dass ein Großteil jener Personen, die von Barrieren betroffen waren und sind, nicht beachtet wurden (etwa die oben aufgeführten Personen mit Lernschwierigkeiten, um nur ein Beispiel zu nennen). Dabei gibt es unzählige verschiedene Arten und Weisen, Barrierefreiheit in das alltägliche Leben zu integrieren – sei es durch die Zugänglichkeit einer Webseite für die unterschiedlichsten Personengruppen oder durch das Bereitstellen von Rampen für Rollstuhlfahrer:innen (vgl. Díaz Cintas et al. 2007: 13).

In dieser Hinsicht können die Konzepte von Barrierefreiheit und Translation sehr gut miteinander verglichen werden, denn auch Translation bietet seit Jahrhunderten jenen Menschen im weitesten Sinne Zugang zu Texten. In der Forschung der letzten Jahre und Jahrzehnte ist vor allem die audiovisuelle Translation in den Fokus gerückt, welche sich mit der Übersetzung audiovisueller Materialien, darunter Filme, Computerspiele, Webseiten und Cartoons beschäftigt. In diesem Feld, der barrierefreien audiovisuellen Translation, gibt es wiederum zwei primäre große Forschungsinteressen, nämlich die SDH-Untertitelung und die Audiodeskription, welche hör- respektive sehbehinderten Menschen den Zugang zu audiovisuellem Material ermöglicht (vgl. Díaz Cintas 2007: 13f.).

Das stetig wachsende Bewusstsein und die Gesetzgebung in Hinblick auf Barrierefreiheit sollen in Zukunft dazu führen, dass barrierefreie Elemente bereits in die Anfangsphase der Produktentwicklung integriert werden. Nur so kann sichergestellt werden, dass Produkte für die unterschiedlichsten Nutzer:innengruppen zugänglich sind – egal, welche Barrieren (körperlich, geistig, sprachlich) dem Zugang im Weg stehen mögen. Hier ist auch wichtig hinzuzufügen, dass das sofortige Miteinbeziehen barrierefreier Inhalte weniger zeit- und kostenintensiv ist als das Hinzufügen im Nachhinein (vgl. Díaz Cintas 2007: 14).

Auch Greco (2018) unterstreicht diese Tatsache und betont außerdem, dass ein reaktiver Zugang zu Barrierefreiheit immer mit gewissen Einschränkungen verbunden ist, welche das Endprodukt meistens entweder weniger funktional oder qualitativ minderwertig machen. Doch das Annehmen eines proaktiven Zugangs bedeutet nicht nur, dass das Umsetzen barrierefreier Elemente von einer späteren in eine frühere Produktdesignphase verlegt wird, sondern dass die Nutzer:innen und zukünftigen Nutzer:innen, aber auch Expert:innen in die Mitte des Designprozesses geholt werden (vgl. Greco 2018: 212f.). Das zeigt, dass ohne Partizipation keine wahre Barrierefreiheit gegeben sein kann, was wiederum die Erkenntnisse aus Kapitel 2.1.2 unterstreicht.

2.2.1 Media Accessibility

Eine Art der Barrierefreiheit, die vor allem auch für die Translationswissenschaft relevant ist und umgekehrt, ist die *Media Accessibility*, also die Barrierefreiheit in Bezug auf Medien, welche ihren Anfang mit SDH-Untertitelungen und Audiodeskriptionen (AD) nahm. Diese ersten Schritte decken sich auch mit der ersten von drei Entwicklungsstufen der Media Accessibility nach Greco (2018). Er beschreibt, dass zu Beginn vor allem auf Menschen mit Behinderungen eingegangen wurde. Nach und nach wurde das Feld von SDH und AD auch auf andere Bereiche, etwa Gebärdensprachdolmetschen, erweitert, jedoch lag der Fokus nach wie vor auf Sinnesbehinderungen. Im zweiten Schritt habe sich der Blick auf Media Accessibility ausgeweitet, wodurch auch sprachliche Barrieren miteinbezogen wurden. Die letzte Entwicklungsstufe, die universalistische, wie Greco (2018) sie nennt, sieht Media Accessibility als eine Möglichkeit, allen Menschen Zugang zu Medienprodukten, Dienstleistungen und Umgebungen zu gewährleisten, die diese in ihrer ursprünglichen Form nicht in Anspruch nehmen könnten (vgl. Greco 2018: 211). Die EU-Richtlinie über audiovisuelle Mediendienste sieht Barrierefreiheit nicht nur als Service für Menschen mit Sinnesbehinderung(en), sondern auch für ältere Personen. Die neueste Untertitelungsnorm ISO/IEC 20071-23:2018 sieht sogar vor, dass barrierefreie Untertitel für eine noch größere Zielgruppe bereitgestellt werden, darunter Menschen mit Sinnesbehinderung(en), Menschen mit kognitiven Einschränkungen, Menschen, die einen Film in einer Nicht-Erstsprache ansehen und Menschen, die die Tonspur nicht abrufen können. Auch Umwelteinflüsse werden berücksichtigt, etwa Situationen, in denen der Ton des audiovisuellen Materials nicht zugänglich (*accessible*), verfügbar (*available*) oder angebracht (*appropriate*) ist. Die neuen gesetzlichen Regelungen zeigen also, dass Barrierefreiheit nicht nur für Menschen mit Sinnesbehinderungen hilfreich sein kann, sondern auch für ältere Menschen, Menschen mit Lernschwierigkeiten und auch Menschen, die Medien in einer Fremdsprache konsumieren (vgl. Romero-Fresco 2019a: 9).

Die traditionelle, enge Sichtweise auf Media Accessibility, welche sich auf Personen mit Sinnesbeeinträchtigungen beschränkt, birgt einige Risiken, darunter auch eine terminologische: Die Bezeichnung „Subtitles for the Deaf and Hard of Hearing“ impliziert schon, dass diese für Menschen mit Hörbeeinträchtigung gedacht sind. Das bedeutet auch, dass andere mögliche Zielgruppen ausgeblendet werden. Neves (2018: 84) schlägt daher einen neuen Begriff vor – *enriched subtitles* –, um die verwirrende und unpassende Terminologie in diesem Bereich und das Stigma rund um das Thema Behinderung zu beseitigen und neue Entwicklungen voranzutreiben.

Darüber hinaus kann dieser Zugang zu Media Accessibility dazu führen, dass ihr Potenzial für sozialen Wandel gehemmt wird. Wenn Barrierefreiheit nur als Service für Personen mit Sinnesbeeinträchtigungen gesehen wird, werden ihre Benefits für alle anderen Zielgruppen – ältere Menschen, Menschen mit Lernschwierigkeiten, sprachliche Minderheiten, Sprachenlerner:innen oder Migrant:innen – außer Acht gelassen. Diesen Umstand bezeichnet Greco (2018: 212) als „ghetto effect“, welcher die Media Accessibility stark behindern kann. Romero-Fresco (2019) nennt hier das Beispiel von Südafrika, wo laut Gesetz alle elf Sprachen den gleichen Status haben sollten und Services in all diesen Sprachen zur Verfügung gestellt werden sollten, aber de facto 76 Prozent der Programme der SABC (South African Broadcasting Corporation) englisch sind. Außerdem stellt die SABC ausschließlich englische Untertitelungen bereit. Obwohl Media Accessibility genau in solchen Fällen zu mehr Sprachengleichheit führen könnte, was wiederum Alphabetisierung, Bildung, sozialen Zusammenhalt und Lebensqualität vorantreiben würde, wird sie lediglich als Hilfsmittel für Menschen mit Sinnesbeeinträchtigungen gesehen und kann so nicht ihr vollstes Potenzial erreichen (vgl. Romero-Fresco 2019a: 10).

2.2.2 Anwendungsbereiche der Media Accessibility

Um die Vielfältigkeit der Media Accessibility zu illustrieren, sollen hier einige Beispiele für Anwendungsbeispiele gegeben werden. Es ist jedoch wichtig zu erwähnen, dass dies keine exhaustive Liste ist und alle genannten Felder aus Platzgründen nur angeschnitten werden können.

Wie bereits erwähnt, nahm die Media Accessibility ihren Anfang in der Erstellung von *Untertiteln für Hörgeschädigte und Gehörlose (SDH)* und *Audiodeskriptionen (AD)*, welche hör- beziehungsweise sehbehinderten Menschen den Zugang zu audiovisuellen Materialien ermöglichen sollte. Da die SDH-Untertitelung als Fokus dieser Masterarbeit im folgenden Kapitel genauer dargestellt wird, erfolgt hier nur eine kurze Skizzierung: Unter Untertitelung für Hörgeschädigte und Gehörlose versteht man neben der bloßen Transkription des Gesprochenen auch das Hinzufügen von Situations- und Kontextangaben (z.B. Speaker-IDs, Sprechmodalitäten) sowie von Geräuschen und Musik (vgl. Jüngst 2020: 213f).

Als Audiodeskription wird die ergänzende Beschreibung eines Films oder einer Veranstaltung bezeichnet. Wichtig ist hierbei, dass diese Beschreibungen in Gesprächspausen eingebracht werden und nur Elemente enthalten sollen, welche sich nicht aus der Audiospur

ergeben. Audiodeskriptionen können nicht für sich selbst stehen, sondern benötigen den Filmton, um ihren Sinn zu erfüllen (vgl. Jüngst 2020: 175f.).

Neben diesen traditionellen Ausprägungen der Media Accessibility haben sich auch andere Felder entwickelt, etwa das *Schriftdolmetschen* – eine Verschriftlichung eines mündlich dargebotenen Ausgangstextes, welche sich nahezu in Echtzeit, in wörtlicher oder komprimierter Form, vollzieht und sich primär an Menschen mit Hörbeeinträchtigung richtet. Ähnlich wie bei der SDH-Untertitelung werden auch hier außersprachliche Merkmale, wie Sprecher:innen, Sprecher:innenwechsel, nonverbale Äußerungen und Geräusche wiedergegeben. Dabei können verschiedene Methoden verwendet werden: Eingabe des Textes über eine Tastatur, Respeaking (Eingabe mittels Spracherkennung) oder auch Stenografie. Die Zielgruppe ist sehr divers und kann von wörtlichen (für Menschen mit hoher Schriftsprachkompetenz) oder komprimierten (bei niedriger Schriftsprachkompetenz) Versionen des Textes Gebrauch machen. Generell kommt das Schriftdolmetschen dort zum Einsatz, wo es wichtig ist, dass die betroffene Person mit hoher Sicherheit alles Gesagte versteht, also etwa bei Arztgesprächen oder in der Schule. Im Moment wird diese Art des Dolmetschens vor allem intralingual ausgeübt; an der didaktischen Entwicklung des interlingualen Schriftdolmetschens an der Roehampton University in London wird jedoch bereits geforscht (vgl. Witzel 2018: 304ff.).

Eine weitere Form der Media Accessibility ist das *Gebärdensprachdolmetschen*, welches als relativ junge Disziplin gilt: Seine Professionalisierung wurde in Deutschland erst in den 1970er-Jahren durchgeführt, obwohl das Bewusstsein um Gebärdensprachen auch davor schon lange vorhanden war. Diese Art des Dolmetschens kann sich auf drei Arten vollziehen – Als Simultan- (beinahe gleichzeitige Übertragung des Gesprächsinhalts) oder Konsektivdolmetschung (Übertragung mit Verzögerung) oder Vom Blatt-Übersetzung (mündliche Übersetzung eines schriftlich dargebotenen Textes) (vgl. Benner & Herrmann 2018: 388f.). Bisher sind die meisten Ausbildungsgänge für hörende Dolmetscher:innen konzipiert, es gibt aber einige Kursangebote, welche auch taube Dolmetscher:innen ausbilden (vgl. Benner & Herrmann 2018: 391).

Darüber hinaus gibt es Übersetzungen in *Leichte Sprache*, einer „Varietät des Deutschen mit reduziertem lexikalischen und grammatischen Inventar“ (Maaß 2018: 273). Diese Varietät setzt auf einen zentralen, alltagsnahen Wortschatz, welcher möglichst neutral, präzise und ohne Nebenbedeutungen ist. Außerdem werden etwa Fachwörter so gut wie möglich vermieden, und Komposita werden entweder durch Bindestrich (Narkose-Facharzt) oder Mediopunkt (Zusammen·arbeit) in ihre Bestandteile aufgeteilt. Da die meisten Leichte-Sprache-Texte keine

Neuschaffungen sind, sondern sich an bestehende Textangebote anlehnen, gelten diese als intralinguale Übersetzungen (vgl. Maaß 2018: 273f.).

Unter dem Begriff *Unterstützte Kommunikation* (englisch: *Augmentative and Alternative Communication*) werden kommunikationsbezogene Hilfen und Hilfsmittel für Personen zusammengefasst, die sich vorübergehend oder dauerhaft nicht oder nur eingeschränkt verbalsprachlich mitteilen können. Im deutschsprachigen Raum hat sich die Bezeichnung „Unterstützte Kommunikation“ als Oberbegriff durchgesetzt, während im Englischen in *Augmentative Communication* (ergänzende und zusätzliche Kommunikationsformen, z.B. Visualisierungen) und *Alternative Communication* (alternative Formen der Kommunikation wie Braille-Schrift) differenziert wird (vgl. Musenberg 2018: 361). Auch diese Zielgruppe ist sehr divers – zu ihr zählen etwa Personen mit Autismus-Spektrum-Störung, verzögerter Sprachentwicklung, Menschen mit Aphasie nach Schlaganfällen oder Schädel-Hirn-Traumata oder mit kognitiven Beeinträchtigungen (vgl. Musenberg 2018: 364).

2.2.3 Wissenschaftliche Verortung

Viele Expert:innen stellen sich also die Frage, zu welchem wissenschaftlichen Feld die Media Accessibility gehören soll. Zu Beginn wurde sie aufgrund ihres Anwendungsfeldes (SDH und AD) als Teilbereich der audiovisuellen Translation gesehen, wodurch sie als translationswissenschaftliches Feld galt. In den letzten Jahrzehnten zeigte sich aber, dass Barrierefreiheit im wissenschaftlichen Bereich einen interdisziplinären Zugang verlangt, um auf alle Teilbereiche eingehen zu können. So ist es auch bei der Media Accessibility, denn sie vereint Ansätze aus audiovisueller Translation, Computertechnologie, Psychologie, Gesundheitswesen, Anthropologie, Kulturwissenschaften und einigen mehr (vgl. O'Hagan & Mangiron 2013: 36f.). Und auch wenn ein bestimmter Teil der Media Accessibility (etwa nur die Audiodeskription) untersucht wird, ist, so Expert:innen, ein interdisziplinärer und internationaler Ansatz am zielführendsten (vgl. Greco 2018: 215).

Das Feld der Media Accessibility war also zu Beginn ein rein translationswissenschaftliches Feld, ein Teilbereich der audiovisuellen Translation. Nach einiger Zeit wurden jedoch auch Personen mit sprachlichen Barrieren miteinbezogen, wodurch sich Media Accessibility und AVT überlappten, sich aber immer noch im Bereich der Translationswissenschaft bewegten. Doch sobald audiovisuelle Problematiken in einem barrierefreien Kontext gesehen werden, verlässt man den translationswissenschaftlichen Raum,

da Translation nur noch einer von vielen Faktoren der Barrierefreiheit ist, und außerdem nicht immer der wichtigste (vgl. Romero-Fresco 2018: 200f.). Das Einbetten der Media Accessibility in das Feld der Translationswissenschaft hat außerdem zur Folge, dass man sich in einem kleineren Rahmen mit begrenzten Modalitäten bewegt, indem man nur jene aus der Translationswissenschaft benutzt und andere, die nichts oder wenig mit Translation zu tun haben, außer Acht lässt. Die Media Accessibility muss also als breiteres, interdisziplinäres Feld gesehen werden, welches verschiedene Felder miteinander vereint (vgl. Greco 2018: 217f.).

Die Entwicklung der Media Accessibility vollzieht sich jedoch nicht ohne Probleme: So kritisiert Romero-Fresco (2015), dass, während die Anzahl und die Qualität von Untertiteln für Hörgeschädigte und Gehörlose exponentiell wachsen, die Forschung zeigt, dass diese für gehörlose Rezipient:innen in den meisten Fällen nicht so gut geeignet sind wie für hörgeschädigte. Außerdem würden ihnen oftmals Inhalte in Gebärdensprache(n) fehlen, und die SDH-Untertitel würden aufgrund ihres Inhalts, ihrer Struktur und ihrer Geschwindigkeit oft zu Verständnisproblemen führen (vgl. Wehrmeyer 2015: 218f.). Dies unterstreichen auch Romero-Fresco & Dangerfield (2022), denn die Obergrenze für Zeichen pro Sekunde in Untertiteln, welche bei Netflix bei 20 CPS liegt, sei das Limit für hörende Rezipient:innen – für gehörlose Personen sind die Untertitel daher oft nicht nutzbar. Er sieht den Grund dafür in der Wirtschaftlichkeit, denn höhere CPS bedeuten auch weniger Anpassungen (z.B. durch Auslassungen) und somit weniger Zeit für die Erstellung der Untertitel. Das Problem dabei ist, dass die Zielgruppe, für die diese Untertitel primär angefertigt wurde, diese nun nicht mehr nutzen kann, während hörende Rezipient:innen sie als zusätzliche Hilfe verwenden können. Viele Forscher:innen plädieren daher für eine Reduzierung der CPS in Untertitelungsrichtlinien – im deutschsprachigen Raum sind zum Beispiel maximal 15 Zeichen pro Sekunde in SDH-Untertiteln erlaubt –, damit gehörlose Personen diese auch nutzen können (vgl. Romero-Fresco & Dangerfield 2022: 19).

Ein weiterer Aspekt, den Romero-Fresco (2018) betont, ist, dass sowohl die Theorie als auch die Praxis der Media Accessibility oft für, aber nicht von Menschen mit Behinderung ausgeführt wird. Dies kann, wie bereits in Kapitel 2.1.2 beschrieben, zu einem paternalistischen und eigennützigen Zugang führen, der die aktive Rolle der *able-bodied*⁴ und die passive Rolle von Menschen mit Behinderung aufrechterhält. Ein fairerer und inklusiverer Zugang zu Media

⁴ Die offizielle deutsche Entsprechung dieses Begriffs – nichtbehindert bzw. Nichtbehinderte:r – wird von vielen Menschen mit Behinderung als diskriminierend empfunden und daher im deutschsprachigen Diskurs oftmals vermieden. Einige argumentieren außerdem, dass der Begriff „nichtbehindert“ nicht das Gegenteil von „behindert“ sei – er zeige nur, dass „die breite Masse noch nicht verstanden habe, dass Behinderungen zum Leben dazugehören“ (Ein Blog von Vielen 2020).

Accessibility sollte daher zwei Hauptaspekte miteinander vereinen: Zugang zu Inhalten (*access to content*) und Zugang zu Kreation (*access to creation*) (vgl. Romero-Fresco 2018: 190).

2.2.4 Access to content

Lange Zeit lag der Fokus der Media Accessibility auf dem *access to content*, also der bloßen Bereitstellung von Inhalten. Romero-Fresco (2018) illustriert anhand des Beispiels des interlingualen Respeaking, inwiefern diese eingeschränkte Sicht auf Media Accessibility dieses Feld beeinflusst und in der Vergangenheit beeinträchtigt hat. Das Respeaking ist eine komplexe, mit dem Simultandolmetschen vergleichbare Aufgabe, welche jedoch in vielen Fällen von Untertitler:innen ausgeführt wird. Dabei wird oft vergessen, dass das Produkt, nämlich die Live-Untertitel, häufig von einer diversen Zielgruppe (mit und ohne Hörbeeinträchtigung) konsumiert wird. Dies hat zur Folge, dass einige Rezipient:innen womöglich Nachteile im Verständnis haben – etwa hat sich gezeigt, dass Menschen mit Gebärdensprache als Erstsprache die Live-Untertitelung oftmals nicht so gut verstehen wie hörgeschädigte Menschen, da diese beiden Gruppen ein unterschiedliches Niveau an Schriftsprachkompetenz aufweisen und die Untertitel oftmals schnell und sehr wörtlich sind. Gleichzeitig können auch hörgeschädigte Rezipient:innen benachteiligt werden, wenn etwa nonverbale Elemente nicht in die Untertitelung einbezogen werden, welche für das Verständnis nötig sind. Für hörende Rezipient:innen jedoch seien diese nonverbalen Elemente nicht störend, weshalb Szarkowska (2013) und Romero-Fresco (2018) sich dafür einsetzen, dass diese bei interlingualem Respeaking automatisch inkludiert werden. Trotz allem müsse dafür gesorgt werden, dass die Qualität der Untertitelung nicht unter dem zusätzlichen Informationsangebot leidet (vgl. Romero-Fresco 2018: 191f.).

Romero-Fresco (2018: 192) folgert, dass das interlinguale Respeaking einer der neuen, aufstrebenden Bereiche der Media Accessibility ist, der dazu beitragen kann, dass dieses Feld einer breiteren Öffentlichkeit zugänglich gemacht wird und diese inklusiver und empathischer macht. Die hörenden Rezipient:innen sollen sich demnach leichter in die Lage gehörloser Menschen versetzen können, wobei gleichzeitig die Verbreitung dieser Untertitelungen nur durch ein großes Zielpublikum sichergestellt werden kann.

2.2.5 Access to creation: Accessible Filmmaking

Das Accessible Filmmaking versteht sich als die Integration von Translation und Barrierefreiheit in den Filmproduktionsprozess. Es kann dabei helfen, die Lücke zwischen

audiovisuellem Produkt und audiovisueller Translation bzw. Media Accessibility zu schließen und die Sichtbarkeit der barrierefreien und translatorischen Services in audiovisuellen Produktionen zu erhöhen. Dabei ist es wichtig, erneut zu betonen, dass alle möglichen Zielgruppen von dieser Zusammenarbeit profitieren können, von Menschen mit Behinderung(en) bis hin zu Sprachenlerner:innen (vgl. Romero-Fresco 2019a: 17).

Um diesen Standard des Filmemachens zu erreichen, müssen sowohl die Inhalte als auch der kreative Prozess zugänglich gemacht werden. Einige Projekte bedienen sich des AFM-Modells, um Teilnehmer:innen mit Sinnesbeeinträchtigung eine Möglichkeit zu geben, audiovisuelle Produkte zu kreieren. Die daraus resultierenden Filme sind oftmals innovativ und einzigartig und werden als Teil des *non-cinema* angesehen. Als *non-cinema* werden Filme bezeichnet, die sich nicht an traditionelle Produktions- und Vertriebswege halten und sich mit nicht oder wenig repräsentierten Aspekten beschäftigt. Auch wird eine Realität gezeigt, welche in Kinos in den meisten Fällen entweder gar nicht oder nur aus zweiter Hand von nicht-behinderten Produzent:innen gezeigt wird (vgl. Romero-Fresco 2019a: 17).

Im Kampf um den access to creation ist es nicht unüblich, einen noch breiteren Zugang zu Media Accessibility zu wählen, welcher auch Diversität fördern soll. So sollen auch andere unterrepräsentierte Gruppen, wie People of Colour oder LGBTQ+, sichtbarer gemacht und in den Produktionsprozess eingebunden werden. Dennoch, betonen Expert:innen, ist es noch ein weiter Weg, bis von wahren access to creation die Rede sein kann. Denn bis dahin müssen viele Herausforderungen überwunden werden, etwa „access to (assistive) technology, access to funding and access to the professional market, which in itself requires a profound change in attitudes and perhaps also a set of proactive policies“ (Romero-Fresco 2019a: 17). Zu guter Letzt sollte auch normalisiert werden, dass behinderte Produzent:innen sich nicht nur mit Themen rund um Behinderung (*disability art*) auseinandersetzen „dürfen“, sondern um verschiedenste Themen (*disability in the arts*). Es wäre außerdem wünschenswert, Filme dieser Art der breiten Masse zugänglich zu machen, anstatt diese als *non-cinema* zu bezeichnen (vgl. Romero-Fresco 2019a: 17).

Wie in Kapitel 3 noch detaillierter beschrieben wird, wurden sowohl die Untertitelung für Hörgeschädigte und Gehörlose als auch die Audiodeskription als translatorische Prozesse von Beginn an ausgelagert. Dies gilt sowohl für die Film- als auch die Fernsehindustrie, da barrierefreie Services schon seit jeher als kostspielig und als nur für eine begrenzte, sehr spezifische Zielgruppe zugeschnitten empfunden wurden. Diese Ansicht manifestiert sich auch in der Tatsache, dass SDH in beiden Branchen von einem Drittanbieter erstellt und verbreitet

wurden und nicht in Zusammenarbeit mit dem Kreativteam der Film- oder Fernsehproduktion erstellt wurden (vgl. Romero-Fresco 2019a: 30).

Die erste Initiative, dieses Kurations- und Distributionsmodell zu ändern, entstand in Anlehnung an das Konzept des *universal design*, welches besagt, dass Gebäude, Produkte und Dienstleistungen für Menschen mit und ohne Behinderung gleichermaßen zugänglich sein soll. Um den Prinzipien des *universal design* zu entsprechen, muss das Design eines Produkts so viele Endnutzer:innen wie möglich inkludieren – und zwar von Anfang an (vgl. Romero-Fresco 2019a: 30f.). Udo & Fels (2010) wandten dieses Konzept auf SDH und Audiodeskription an, um festzustellen, ob diese Angebote als *universal design* gelten könnten. SDH werden von vielen Expert:innen als sogenannte *electronic curb-cuts* bezeichnet; Das sind Produkte, die nicht nur den Zielbenutzer:innen, sondern auch weniger mitgedachten Nutzer:innen Vorteile bringen. Bei SDH sind das nicht nur hörgeschädigte und taube Menschen, sondern etwa auch Sprachenlerner:innen. Udo & Fels (2010) folgern dennoch, dass die beiden Angebote keine Beispiele für *universal design* sein können, da sie nicht von Anfang an in den Produktionsprozess eingegliedert sind. Da sie jedoch die Interpretation des Publikums stark beeinflussen können, schlagen die Autor:innen folgenden Ablauf für SDH und Audiodeskription vor:

We assert that audio describers and captionists should operate under a similar system [to the rest of the filmmaking crew], reporting to or, at least, consulting with a director of accessibility services. This team would then meet with the production's director to develop an accessibility strategy that re-interprets the "look and feel" of the production. The captioning and description team would then work together to develop prototypes that would, in turn, be approved by the director before being produced. The final product should receive similar attention. (Udo & Fels 2010: 218)

Genau diese Strategie wurde 1997 von der sehbehinderten britischen Filmproduzentin Raina Haig für den Film „Drive“ angewendet. Dieser Film war der erste, der Audiodeskription von Anfang an in den Produktionsprozess eingliederte. Für Haig (2002) ist die Audiodeskription immer in Rücksprache oder gar in Zusammenarbeit mit den Produzent:innen der Filme zu erstellen, da nur so sichergestellt werden kann, dass auch sehbehinderte Menschen einen angemessenen Zugang zum dargebotenen audiovisuellen Material haben und auch die künstlerische Gestaltung im Sinne der Produzent:innen ist (vgl. Haig 2002: o.S.). Das bedeutet jedoch nicht, dass die Produzent:innen auch Expert:innen im Gebiet der barrierefreien Services sein müssen. Sie können sich, genauso wie bei anderen Prozessen der Filmproduktion, auf die Dienstleister:innen verlassen, sollten jedoch über Grundlagenwissen zu Audiodeskription verfügen, um angemessene Entscheidungen treffen zu können (vgl. Udo & Fels 2010: 218).

Haig (2002: o.S.) betont, dass im Gegenzug auch Audiodeskripteur:innen sich der Vision der Produzent:innen anpassen können müssen, um etwa künstlerische Entscheidungen zu verstehen.

Obwohl das Konzept des universal design eine Schlüsselrolle für das Accessible Filmmaking spielte, kritisieren einige Forscher:innen, dass dieses Konzept nicht ganz mit SDH und Audiodeskription vereinbar ist. Vor allem die Tatsache, dass Udo & Fels (2010) nur barrierefreie, nicht aber translatorische Services, in ihre Analyse miteinbezogen, wurde in Expert:innenkreisen immer wieder als problematisch betrachtet, da ein Modell so weitreichend und kosteneffizient wie möglich sein sollte, um von der Industrie angenommen zu werden. Wenn es sich jedoch nur auf barrierefreie Services bezieht, besteht das Risiko, dass das Modell als zu kostspielig und die Zielgruppe als zu spezifisch angesehen wird. Das Accessible Filmmaking versucht daher alle Aspekte zu vereinen, die benötigt werden, um einen Film etwa seh- und hörbehinderten Menschen, aber auch Sprachenlerner:innen zugänglich zu machen, und spricht damit eine große Zielgruppe an. Außerdem soll betont werden, dass die Veränderungen im Produktionsprozess das Filmerleben der ursprünglichen Zielgruppe oder die kreative Vorstellung der Produzent:innen nicht verändern sollen. Vielmehr soll diese Vision auch in der Translation bestehen, und Produzent:innen sollen selbst entscheiden, inwiefern sie mit dem translatorischen Team zusammenarbeiten (vgl. Romero-Fresco 2013).

3. Audiovisuelle Translation und (SDH-)Untertitelung im Wandel der Zeit

3.1 Audiovisuelle Translation

Die audiovisuelle Translation (AVT) hat sich in den letzten Jahrzehnten zu einem der wichtigsten Forschungsbereiche innerhalb der Translationswissenschaft entwickelt. Dabei versteht man unter dem Begriff *audiovisuelle Translation* das „Übersetzen von Medienformaten, die einen sichtbaren und einen hörbaren Teil haben“ (Jüngst 2010: 1), weshalb sie als Form der multimodalen Kommunikation gelten. Das vorliegende Material wird zum Teil verändert, zum Teil bleibt es gleich oder wird durch neue Elemente ergänzt. Zur audiovisuellen Translation zählen unter anderem folgende Bereiche:

- die Synchronisation, bei der die ausgangssprachliche Dialogspur durch eine zielsprachliche ersetzt wird;
- die Untertitelung, die in inter- und intralinguale Untertitelung unterteilt werden kann. Bei interlingualen Untertiteln wird ausgangssprachliche Dialogspur in die Zielsprache übersetzt und in Untertitel umgewandelt, sodass das Zielpublikum zwar die ausgangssprachliche Dialogspur hört, aber zielsprachliche Untertitel liest. Im Gegensatz dazu sind bei intralingualen Untertiteln die Dialogspur und die Untertitel in der gleichen Sprache abgefasst;
- Hörfilme bzw. Filme mit Audiodeskription, welche für blinde und sehbehinderte Menschen gedacht sind und die auch Bildinhalte akustisch beschreiben;
- und die Voice-over-Übersetzung, bei der man neben der zielsprachlichen Dialogspur auch die ausgangssprachliche Dialogspur im Hintergrund hört (vgl. Jüngst 2010: 1-3).

Intralinguale Untertitel können wiederum in zwei Gruppen mit zwei unterschiedlichen Zwecken unterteilt werden. Bei der ersten Kategorie handelt es sich um intralinguale Untertitel, zum Beispiel zum Zweck des Sprachenlernens für Hörende – hier werden keine Hintergrundinformationen, wie Geräusche oder Musik, gegeben, und der gesprochene Text wird wörtlich übernommen. Die zweite Kategorie von intralingualen Untertiteln ist die *Untertitelung für Hörgeschädigte und Gehörlose*, die SDH-Untertitelung, die sowohl verbale als auch nonverbale Elemente der untertitelten Szene enthält (vgl. Gambier 2013: 49). Diese Art der Untertitelung wird als Hauptaugenmerk der vorliegenden Masterarbeit im weiteren Verlauf dieses Kapitels noch genauer beschrieben.

3.1.1 Historische Entwicklung der AVT

Obwohl die Praxis der audiovisuellen Translation seit der Erfindung des Films und des Kinos Ende des 19. Jahrhunderts existierte, wurde dieses Feld, wie bereits erwähnt, bis in die 1990er-Jahre nicht in der Wissenschaft diskutiert (vgl. Díaz Cintas & Remael 2021: 1). Dabei war das Übersetzen schon von Beginn an eine der größten Herausforderung in der Filmgeschichte, denn auch in der Ära des Stummfilms waren die Filme alles andere als stumm. So wurden etwa von den Filmerzähler:innen oder -erklärer:innen beispielsweise Geschichten erzählt oder Zwischentitel vorgelesen, auch begleitende Musik oder Soundeffekte wurden abgespielt (vgl. Gambier 2013: 45). Die Filmerzähler:innen oder -erklärer:innen wurden in den USA und Europa bis in die 1910er-Jahre vor allem dafür eingesetzt, die eingeblendeten Zwischentitel für Analphabet:innen vorzulesen oder zu erklären. Darüber hinaus übersetzten sie auch fremdsprachige Inhalte der Filme, weshalb ihre Tätigkeit heute als intralinguale, interlinguale oder auch intersemiotische Translation bezeichnet würde. Einige Jahre später, in den 1920er-Jahren, wurden die ersten *Talkies* entwickelt – die ersten Filme mit (vorerst geringem) Sprechanteil. Der erste Film dieser Art war „The Jazz Singer“, veröffentlicht 1927 von Warner Brothers. Es ist nicht genau überliefert, wie dieser Film in nicht-englischsprachigen Kinos gezeigt wurde und ob bzw. wie er übersetzt wurde. In Frankreich wurde er demnach mit französischen Zwischentiteln gezeigt, und die gesprochenen Sequenzen wurden ebenfalls übersetzt (vgl. O’Sullivan & Cornu 2018: 17).

Ungefähr im gleichen Zeitraum begannen die Filmproduktionsfirmen, mit Tonspuren zu arbeiten, was auch Änderungen in der Produktion und in der Bildeinstellung mit sich trug. Mit dieser Neuerung konnten die Produzent:innen mehrere Sprachfassungen desselben Films aufnehmen und diese besser an das Zielpublikum anpassen. Zu Beginn wurden noch Schauspieler:innen bzw. Sprecher:innen beispielsweise aus Deutschland oder Frankreich nach Hollywood eingeflogen, um die verschiedenen Fassungen aufzunehmen, bevor – vor allem aus finanziellen Gründen – schließlich auch die Produktion in die Zielländer verlagert wurde. Als in den 1930ern auch die Synchronisation entwickelt wurde, wurde die Verantwortlichkeit für die sprachlichen Aspekte auch an spezielle Synchronisationsfirmen weitergegeben und damit gänzlich ausgelagert. Auch Remakes waren eine Form des Umgangs mit der Sprachenvielfalt, wobei hier einige Aspekte des Originalfilms beibehalten und andere teilweise verändert werden. So wird etwa die Sprachfassung neu erstellt, während Teile des Inhalts gleichbleiben oder variieren können. Während zu Beginn dieser Technik vor allem Remakes von amerikanischen Filmen in Europa angefertigt wurden, so hat sich dieser Trend seit den 1980er-

Jahren umgekehrt. Darüber hinaus werden schon seit jeher multilinguale Filme produziert, welche Sprachidentität(en), -kontakte und -konflikte abzubilden versuchen (vgl. Gambier 2013: 46).

3.1.2 Untertitelungs- vs. Synchronisationsländer

Im Bereich der audiovisuellen Translation bildeten sich bereits in den 1930er-Jahren die Lager Synchronisation vs. Untertitelung. Es ist nicht klar zu definieren, welche genauen Faktoren für die Wahl des einen oder anderen Translationsprozesses entscheidend waren; zumeist waren es verschiedene wirtschaftliche, ideologische oder pragmatische Aspekte. In Frankreich wurden bis in die 1950er-Jahre beispielsweise beide Formen immer wieder angewendet, was auch durch den langen Konkurrenzkampf zwischen französischen und US-amerikanischen Filmemacher:innen begründet werden kann. In den meisten Fällen ist es jedoch so, dass Synchronisationsländer eher größere, „internationalere“ Sprachgemeinschaften haben, also zum Beispiel Englisch oder Französisch sprechen (vgl. Gambier 2013: 46). So zählen unter anderem Deutschland, Frankreich, Spanien und Italien zu den Synchronisationsländern. Die Menge der Zuschauer:innen, die die zielsprachliche Synchronfassung ansehen werden, ist ausschlaggebend für die Erstellung einer solchen Synchronfassung, da diese teuer und zeitintensiv sind und sich, so Jüngst (2010), erst ab 40 Filmvorstellungen in Kinos lohnen. Im Vergleich dazu sind Untertitelungen günstig und weniger zeitintensiv, weshalb kleinere Sprachgemeinschaften, wie die skandinavischen Länder, Belgien, die Niederlande, Luxemburg und Portugal auf diese zurückgreifen. Gerade bei Ländern mit mehreren Landessprachen bietet sich die Untertitelung an, da man entweder zwei eigenständige Sprachfassungen erstellen oder doppelte Untertitel verwenden kann. Bisher wurde noch kein Wechsel von Synchronisation auf Untertitelung oder umgekehrt festgestellt, die Tradition des einen oder anderen Übersetzungsprozesses dürfte sich wohl auch weiterhin durchsetzen. Denn auch das Publikum von Synchronisationsländern hat sich laut europaweiten Umfragen im Jahr 2006 schon so sehr an die vorherrschenden Kulturen gewöhnt, dass es lieber synchronisierte Filme ansieht als Originale mit Untertiteln. Die englischsprachigen Länder können weder als Untertitelungs- noch als Synchronisationsländer bezeichnet werden, da ohnehin nur sehr wenige fremdsprachige Filme diese Länder erreichen und diese im Normalfall untertitelt werden (vgl. Jüngst 2010: 4-6). Die technologischen Entwicklungen der letzten Jahre und Jahrzehnte (etwa durch Streamingdienste oder ein erweitertes Angebot von barrierefreien Inhalten) machen die Einteilung in Synchronisations- bzw. Untertitelungsländer jedoch zunehmend schwieriger (vgl.

Gambier 2013: 46), da etwa der Streamingdienst Netflix für beinahe alle Filme und Serien sowohl Untertitel als auch zielsprachliche Synchronfassungen bereitstellt.

3.1.3 Eingang in die Translationswissenschaft

Die audiovisuelle Translation wurde in der Translationswissenschaft bis in die 1990er-Jahre aus verschiedenen Gründen nicht als „richtige Übersetzung“ verstanden – unter anderem, weil zeitliche und räumliche Einschränkungen in den Augen vieler Translationswissenschaftler:innen das Endprodukt zu sehr begrenzen würden. Aus diesem Grund sprach man damals eher von Adaptation als von Translation. Außerdem wurde die AVT lange aus translationswissenschaftlicher Sicht außer Acht gelassen, da die Begriffe *Translation* bzw. *Übersetzung* eher für die Übertragung (gedruckter) wörtlicher Inhalte verwendet wurde, nicht aber für andere, multimodale Elemente, wie zum Beispiel Filme (vgl. Díaz Cintas & Remael 2021: 3). Denn genau diese außersprachlichen Aspekte sind es, die zum Beispiel laut O’Sullivan (vgl. 2013: 2) das Verständnis für die audiovisuelle Translation ausmachen – nicht der reine Fokus auf die sprachlichen Faktoren.

Audiovisuelle Texte beruhen auf dem Zusammenspiel mehrerer kommunikativer Codes, etwa Bild und Ton, und gehören daher zur multimodalen Kommunikation. Anders als in der Literatur wird die geschaffene Realität nicht nur durch Wörter, sondern zusätzlich durch bestimmte Bildsequenzen repräsentiert. Folglich vollzieht sich die Vermittlung von Bedeutung sowohl durch den Dialog zwischen zwei Charakteren als auch durch Bilder, Gesten, Kamerabewegungen, Special Effects und weitere filmische Aspekte. Die übertragenen Informationen werden also gleichzeitig über den auditiven und visuellen Kanal aufgenommen. Dieses Zusammenspiel der einzelnen Codes stellt auch die Herausforderung der audiovisuellen Translation dar, da die Translator:innen zumeist nur mit dem sprachlichen Code (Dialog, Voiceover und geschriebene Texte) arbeiten können. Eine weitere Schwierigkeit besteht darin, trotz der Einschränkungen in Bezug auf Bild, Ton und Zeit Synchronität zwischen Ausgangs- und Zieltext herzustellen, denn die Aussage des Ausgangstext soll so synchron und vollständig wie möglich in den Zieltext übertragen werden. Genau diese Einschränkungen waren ein ausschlaggebender Grund dafür, dass die audiovisuelle Translation lange nicht wissenschaftlich untersucht wurde (vgl. Díaz Cintas & Remael 2021: 4).

Eine der ersten Übersetzungswissenschaftler:innen, die die Relevanz der verschiedenen Texttypen für die Übersetzung erkannte und auch auf multimodale Texte hinwies, war Katharina Reiß (1971). Sie unterschied zunächst die drei Texttypen informativ, operativ und

expressiv, bevor sie eine Mischform einführte, welche als *audiomedialer Text* bezeichnet wird. Diese beschreibt jene Texte, die auf nicht-sprachliche (technische) Medien und auf grafische, akustische oder visuelle Ausdrücke angewiesen sind und nur in der Kombination ihr volles Potenzial erreichen (vgl. Reiß 1971: 43). Es handelt sich demnach um Texte, die geschrieben werden, um gesprochen, gesungen oder gehört zu werden, aber nicht, um gelesen zu werden (vgl. Reiß 1971: 27). Einige Jahre später wird dieser Texttyp von Reiß & Vermeer (1984) als *multimediale Texte* bezeichnet und inkludiert nun sowohl auditiv als auch visuell rezipierbare Texte (vgl. Reiß & Vermeer 1984: 211). Als Weiterentwicklung von Reiß' Idee stellte Snell-Hornby (2006) vier weitere Textklassen vor, welche auch alle durch nicht-sprachliche Elemente charakterisiert werden. Multimediale Texte (im Englischen auch *audiovisual texts*) zeichnen sich in diesem Modell dadurch aus, dass sie durch technische und/oder elektronische Medien visuell und auditiv vermittelt werden (z.B. Filme, Über- und Untertitel). Multimodale Texte vereinen Modi verbaler und nonverbaler Kommunikation, welche visuell und auditiv vermittelt wird – hierzu zählen unter anderem Opern. Multisemiotische Texte wiederum bedienen sich verschiedener grafischer verbaler und nonverbaler Zeichensysteme, also beispielsweise Comics, und audiomediale Texte sind jene Texte, die geschrieben werden, um gesprochen zu werden (z.B. Reden). All jene Texte haben gemeinsam, dass sie über die geschriebene Sprache hinausgehen – ein weiterer Grund, weshalb sie lange Zeit nicht aus translationswissenschaftlicher Sicht beleuchtet wurden (vgl. Snell-Hornby 2006: 85).

3.1.4 Terminologische Diskussion

Nicht nur die Klassifikation der in der audiovisuellen Translation behandelten Texte wird in der Translationswissenschaft diskutiert, sondern auch die generelle Terminologie des Feldes. In frühen Werken der 1980er- und 1990er-Jahre wurde die Tätigkeit als *constrained* (Titford 1982; Mayoral Asensio et al. 1988) oder *subordinate translation* (Díaz Cintas 1998) bezeichnet – diese Begriffe wurden jedoch bald verworfen, da sie sehr restriktiv konnotiert waren. Auch die Termini *film translation* (Snell-Hornby 1988) oder *film and cinema translation* (Hurtado 1994) wurden verwendet, um das Anwendungsfeld zu vergrößern, jedoch wurden diese Bezeichnungen bald durch das Aufkommen anderer Medien, z.B. VHS und DVDs als unpassend empfunden. Um die verschiedenen Aspekte der AVT zusammenzuführen, wurden die Bezeichnungen *multimedia translation* (Gambier & Gottlieb 2001), *multimodal translation* (Remael 2001) und *multidimensional translation* (Gerzymisch-Arbogast & Nauert 2005) eingeführt. Diese terminologische Vielfalt gibt Aufschluss darüber, wie sich die Disziplin über

die Jahre verändert und immer wieder versucht hat, neue, aufkommende Konzepte und Technologien einzubeziehen (vgl. Díaz Cintas & Remael 2021: 5f.).

Im heutigen wissenschaftlichen Sprachgebrauch hat sich, auch sprachübergreifend, der Terminus *audiovisuelle Translation* durchgesetzt. Zu Beginn wurde dieser Begriff nur für interlinguale Übersetzungen von einer Ausgangs- in eine Zielsprache verwendet, doch im Laufe der Zeit wurden auch intralinguale Translationen, wie SDH-Untertitelungen und Audiodeskriptionen, damit beschrieben. Diese Entwicklung zeigt, dass der Fokus der audiovisuellen Translation eher auf der Zugänglichkeit und Barrierefreiheit anstatt auf Konzepten von Übersetzung und Nicht-Übersetzung liegt. Es ist in dieser Hinsicht also weniger relevant, dass zwischen zwei Sprachen übersetzt wird, als dass Barrierefreiheit (*accessibility*) für alle möglichen Zielgruppen hergestellt wird (vgl. Díaz Cintas & Remael 2021: 6f.). Diese Barrierefreiheit kann sich laut Gambier (2013) auf fünf verschiedene Arten bemerkbar machen: Der Aspekt der *acceptability* (Annehmbarkeit) bezieht sich hierbei etwa auf Sprachnormen, stilistische Entscheidungen und Terminologie, während *legibility* (Leserlichkeit) die Schriftart, die Position und Anzahl der Untertitel meint. *Readability* (Lesbarkeit) bezieht sich noch spezifischer auf den Aspekt des Lesens, unter anderem in Bezug auf Lesegeschwindigkeiten, -gewohnheiten und -kompetenzen. Die Synchronität (*synchronicity*) beschreibt die Angemessenheit verschiedener miteinander verbundener Aspekte, etwa von Lippensynchronität und Aussage oder auch von Gesagtem und Gezeigtem. Der letzte Punkt, *relevance* (Relevanz), bezieht sich auf jene Informationen, die vermittelt, ausgelassen oder hinzugefügt werden müssen, um die kognitive Anstrengung beim Zusehen oder Zuhören nicht zu erhöhen (vgl. Gambier 2013: 56).

3.2 (SDH-)Untertitelung

3.2.1 Definitorische Annäherung

Die Untertitelung für Hörgeschädigte und Gehörlose ist, wie im vorigen Abschnitt beschrieben, eine Form der intralingualen audiovisuellen Translation, da die Ausgangs- und Zielsprache ident sind. Sie hat vor allem den Zweck, audiovisuelle Programme für Menschen mit Hörbeeinträchtigung zugänglich zu machen, da der Großteil der audiovisuellen Programme auch heute noch nur für „physically able people“ (Díaz Cintas et al. 2007: 11) zugänglich ist (vgl. Díaz Cintas et al. 2007: 11). Neben der „bloßen“ Übertragung der Dialogspur, welche an sich schon aufgrund zeitlicher und räumlicher Einschränkungen eine Herausforderung darstellen kann, müssen auch andere Angaben, wie zum Beispiel

Sprecher:innenidentifikationen, Musik oder andere Geräusche, in den Untertiteln inkludiert werden. Diese Form der Untertitelung wird im Englischen – und mittlerweile auch im Deutschen – zumeist als SDH-Untertitelung (*Subtitling for the d/Deaf and hard of hearing*) bezeichnet. Der vor allem im amerikanischen Sprachraum übliche Begriff *Closed Captions* wird im deutschen Sprachraum heute zunehmend seltener verwendet (vgl. Jüngst 2010: 123).

„Subtitling for the d/Deaf and hard of hearing“ hat sich auch im akademischen Umfeld als Standardterminus etabliert – möglicherweise, weil die Forschung in diesem Bereich hauptsächlich in Europa stattfindet, so Neves (2018). Der Begriff vereint die Diversität der Rezipient:innengruppen, sowohl Hörgeschädigte (*hard of hearing*) als auch Gehörlose (*d/Deaf*). Er wurde auch von zwei großen europaweiten Projekten als Terminologie verwendet – Digital Television for All (DVT4ALL) und Hybrid Broadcast and Broadband TV for All (HBBTV4ALL) – und förderte dadurch das Verständnis, dass durch diesen Service eine breit gefächerte Zuseherschaft erreicht werden kann (vgl. Neves 2018: 83).

Closed Captions sind das Gegenteil der *Open Captions*, welche in die Filmdatei eingebrannt sind und nicht zu- und weggeschaltet werden können, zum Beispiel fremdsprachige Untertitel in einem mehrsprachigen Film, die zum besseren Verständnis übersetzt wurden. Closed Captions können hingegen von den Zuseher:innen selbst zu- und weggeschaltet werden (vgl. Diaz Cintas & Remael 2013: 279). *Forced Narratives* wiederum kommen vor allem in interlingualen audiovisuellen Produktionen vor und sind, wie Open Captions, in die Datei eingebrannt. Sie stellen den Zuschauer:innen für die Handlung relevante Informationen zur Verfügung, welche in einer anderen Sprache am Bildschirm präsent sind, zum Beispiel Übersetzungen von Straßenschildern. In interlingualen Untertiteln zeichnen sich diese zumeist dadurch aus, dass sie gänzlich großgeschrieben werden – eine Ausnahme stellen beispielsweise längere zusammenhängende Texte dar, da die Großschreibung hier die Lesbarkeit stören könnte (vgl. Diaz Cintas & Remael 2021: 26).

Unabhängig von der sprachlichen Ebene der Untertitel gibt es weitere Unterscheidungsmerkmale verschiedener Untertitelungstypen, so zum Beispiel die Vorbereitungszeit. *Vorgefertigte Untertitel* werden vor der Ausstrahlung erstellt, während *Live-Untertitel*, wie der Name schon sagt, erst während der Ausstrahlung erstellt bzw. dazugeschaltet werden. Das *Respeaking* mithilfe von Spracherkennungssystemen hat sich als eines der bewährten Mittel zur Erstellung von Live-Untertiteln etabliert. *Semi-Live-Untertitel* sind eine Mischform der beiden Untertiteltypen, bei der ein:e Untertitler:in vor der Ausstrahlung ein Set von Untertiteln ohne Timecode erstellt und diese während der Programmausstrahlung live zuschaltet (vgl. Diaz Cintas & Remael 2021: 22-23). Die Untertitel können des Weiteren im

Ganzen und auf einmal gezeigt werden – in diesem Fall nennt man sie *Pop-up-Untertitel* –, oder sie erscheinen als *kumulative Untertitel*, bei dem ein Teil der Untertitel vor dem anderen erscheint. Pop-up-Untertitel gelten als einfacher zu lesen und weniger störend als kumulative Untertitel. Diese werden vor allem verwendet, um einer Aussage keine wichtigen Informationen vorwegzunehmen und sind daher insbesondere für Witze und Ratespiele oder bei besonders dramatischen Aussagen geeignet (vgl. Diaz Cintas & Remael 2021: 25).

Untertitel erscheinen in den meisten Fällen horizontal am unteren Bildrand, wobei sie im asiatischen Raum auch vertikal auf beiden Seiten des Bildschirms erstellt werden können. Es ist außerdem möglich, die Untertitel an den oberen Bildrand zu verschieben, wenn sie ansonsten wichtige Informationen, wie Einblendungen oder Forced Narratives, verdecken würden (vgl. O’Sullivan & Cornu 2018: 18).

3.2.2 Besonderheiten der SDH-Untertitelung

Wie bereits erwähnt, verfügen SDH-Untertitel über einige Besonderheiten. Dazu zählen unter anderem die Sprecher:innenidentifikation sowie Musik- und Geräuschuntertitel. Die Sprecher:innenidentifikation dient dazu, dass Rezipient:innen feststellen können, welcher Figur ein ganzer Untertitel bzw. eine Teil des Untertitels zugeordnet ist. Hierzu gibt es verschiedene Möglichkeiten – eine von ihnen ist die Farbvergabe, die jedoch auch einige Diskussionspunkte mit sich bringt. Farben können störend wirken, wenn sie Hierarchiekonventionen folgen, die für die Rezipient:innen (noch) nicht nachvollziehbar sind, wenn beispielsweise die Identität einzelner Figuren noch nicht aufgedeckt werden soll. Außerdem kann die Farbwahl falsche Assoziationen in Bezug auf Gendermerkmale auslösen, wenn beispielsweise ein:e blau eingefärbt:e Sprecher:in automatisch als männlich „gelesen“ wird. Auch die Leserlichkeit leidet bei eingefärbten Untertiteln stark gegenüber weiß auf schwarz oder schwarz auf weiß. Sprecher:innenwechsel werden in den meisten Fällen mithilfe von Bindestrichen, mit oder ohne Leerzeichen vor der Aussage, angezeigt (vgl. Mälzer & Wünsche 2018: 337f.).

Doch nicht nur die Sprecher:innenidentifikation erfolgt in den Untertiteln, auch nonverbale oder paraverbale Eigenschaften der Sprecher:innen können beschrieben werden. Oft sind Rückschlüsse auf diese Modalitäten durch das Absehen am Bildschirm möglich, vor allem bei Gesichtsausdrücken – dies ist jedoch nicht immer der Fall, etwa, wenn die Person nicht zu sehen ist. Zu diesen Merkmalen zählen zum Beispiel Rhythmus, Tempo, Lautstärke, Dialekte, Fremdsprachen und viele weitere Besonderheiten. Diese werden oft auf die gleiche Art untitled wie Musik und Geräusche, also entweder in eckigen oder runden Klammern oder

mithilfe von Asterisken, und befinden sich am Anfang des Untertitels vor dem Gesagten. Ein solcher Untertitel könnte wie folgt aussehen: (leise:) Geh weg! (vgl. Jüngst 2010: 133f.).

Ein weiterer wichtiger Aspekt der SDH-Untertitelung sind Musik- und Geräuschuntertitel, welche auf den Kontext der Aussage bzw. Szene verweisen. Geräusche geben Hörenden Aufschluss darüber, was um sie herum geschieht – in Filmen werden sie bewusst eingesetzt. Bei der Erstellung von SDH-Untertiteln ist daher zu entscheiden, welche Geräusche handlungsrelevant sind und ob diese im Bild erkennbar sind. Von diesen Faktoren hängt ab, ob sie untertitelt werden sollen oder nicht. Die Gestaltungsformen können vielfältig sein: Runde oder eckige Klammern oder auch Asteriske können eingesetzt werden. Im deutschen Sprachraum werden Nominalisierungen, wie * Klopfen *, meist gegenüber ganzen Sätzen (* Es klopft. *) bevorzugt, und Geräusche sollen generell nur dann untertitelt werden, wenn sie bzw. ihre Quelle nicht „sichtbar“ sind. Auch die Untertitelung von Musik kann unterschiedlich realisiert werden. Hier ist ebenfalls zu entscheiden, wie handlungsrelevant die Musik ist, und demnach werden die Untertitel gestaltet. Ist sie sehr relevant, werden Titel und Interpret:in sowie teilweise eine Beschreibung der Musik angegeben; handelt es sich eher um Hintergrundmusik, so reicht meistens eine knappe Beschreibung, etwa * kraftvolle Musik *, aus. Außerdem können Songtexte durch eine Note oder eine Raute, meist am Anfang und am Ende eines Textes, angezeigt werden (vgl. Jüngst 2010: 134ff.).

3.2.3 Die Heterogenität der Zielgruppe

Eine der Besonderheiten der SDH-Untertitelung ist die Diversität der hörgeschädigten Zielgruppe, welche auch mit unterschiedlichen Bedürfnissen bei Untertitelungen bzw. Dolmetschungen einhergeht. Die Problematik liegt darin, dass innerhalb der Zielgruppe gänzlich unterschiedliche Bedingungen herrschen können. So lernen *postlingual Ertaubte*, also jene, die spät ihr Gehör verloren haben oder an Altersschwerhörigkeit leiden, im Normalfall keine Gebärdensprache mehr. Das liegt daran, dass die notwendige Feinmotorik fehlt und das Sprachenlernen im Alter schwieriger ist als in jungen Jahren. Von Geburt an Gehörlose (*prälingual Ertaubte*) hingegen erlernen die Gebärdensprache oftmals als Muttersprache, und erst später oder parallel dazu lernen sie das Lesen und das Lippenlesen (*Absehen*). Da die Gebärdensprache eine andere Grammatik als die Lautsprache hat, sind Untertitel, die die Lautsprache wiedergeben, für prälingual Ertaubte folglich fremdsprachlich. Doch auch in dieser Hinsicht unterscheidet sich diese Zielgruppe, zum Beispiel aufgrund von

unterschiedlicher sprachlicher Ausbildung. In der einschlägigen Literatur wird oftmals erwähnt, dass prälingual Ertaubte über eine schlechte Lesekompetenz verfügen (vgl. Jüngst 2010: 124f.).

Die Zielgruppe der SDH-Untertitel kann grob in drei Gruppen aufgeteilt werden: Schwerhörige, Gehörlose und Ertaubte. Die Gruppe der Schwerhörigen kann wiederum im leichtgradig, mittelgradig und hochgradig Schwerhörige eingeteilt werden; so sind leicht- und mittelgradig Schwerhörige mithilfe von Hörgeräten noch fähig zu hören. Bei Gehörlosen und Ertaubten wird zusätzlich zwischen Menschen mit oder ohne Cochlea-Implantat (eine Hörprothese, welche auditive Sprachwahrnehmung ermöglicht) unterschieden. Wie bereits erwähnt, verfügen all diese Gruppen über unterschiedliche Lesekompetenzen, welche nicht nur Resultat des mangelnden Hörvermögens sein können, sondern auch auf individuell besserer oder schlechterer Sprach- und Lesekompetenz basieren können, genau wie bei hörenden Menschen. Bei Personen, die die Lautsprache als Fremdsprache erlernt haben, ist die Lernanstrengung beim Lesen höher. Die Unterschiede in Bezug auf Lesekompetenz und Resthörvermögen führen dazu, dass Untertitel unterschiedlich gut verstanden werden und dass unterschiedliche Zwecke durch Untertitel erfüllt werden. Personen mit Gebärdensprache als Erstsprache verwenden Untertitel mitunter auch, um die eigenen Kenntnisse der Laut- und Schriftsprache zu verbessern. Feststeht jedoch, dass alle drei Zielgruppen auf Untertitel angewiesen sind, um auf das Fernsehangebot zugreifen zu können (vgl. Jüngst 2010: 125f.). Das bedeutet folglich, dass Untertitler:innen gleichzeitig für zwei Kulturen arbeiten müssen – die Gehörlosen- und die Hörendenkultur. Die Arbeit für diese Zielgruppe sollte, so Neves (vgl. 2009: 152), immer auf Augenhöhe geschehen, und die Veränderungen sollten zugunsten der Zuschauerschaft passieren, die nicht als eine Minderheit angesehen werden soll, sondern als einer von vielen Teilen einer fragmentierten Wirklichkeit.

Viele Hörende haben die Vorstellung, dass Menschen mit eingeschränktem Hörvermögen Aussagen „einfach“ absehen könnten, in der Alltagssprache wird dies als „Lippenlesen“ bezeichnet. Absehen ist jedoch eine schwierige Tätigkeit, die häufig zu Missverständnissen führen kann. Spätertaubte Menschen haben durch ihre Vertrautheit mit der Lautsprache in dieser Hinsicht einen gewissen Vorteil gegenüber prälingual Ertaubten. Nur ein kleiner Teil der deutschen Laute kann gut abgesehen werden – Schätzungen gehen von etwa 30 Prozent aus. In Alltagssituationen müssen gehörlose Menschen oft absehen, wenn sie mit Menschen in Kontakt treten, die keine Gebärdensprache beherrschen. Bei Fremdsprachen ist das Absehen praktisch unmöglich, auch wenn sehr gute Sprachkenntnisse vorhanden sind. Das Absehen von Filmdialogen ist demnach weitaus komplizierter als in der Vorstellung vieler hörender Menschen. Viele Gehörlose berichten außerdem, dass nicht vorhandene

Synchronizität zwischen den Lippenbewegungen und den Untertiteln Misstrauen und Unzufriedenheit hervorruft, weshalb SDH-Untertitler:innen also nicht auf diesen Aspekt verzichten dürfen (vgl. Jüngst 2010: 126f.).

Gehörlose verstehen sich als Angehörige einer eigenen Kultur und haben unter anderem eigene Freizeitvereine, wie Sportverbände oder Theatergruppen. Menschen mit eingeschränktem Hörvermögen stoßen bei Hörenden oft auf Unverständnis – es können Missverständnisse oder Reibungen auftreten, weshalb sie sich unter anderen Gehörlosen oft wohler fühlen als unter Hörenden (vgl. Jüngst 2010: 125). Die Diskriminierung nicht-hörender Menschen wird als *Audismus* bezeichnet und äußert sich durch eine bevormundende oder negative Haltung Hörender gegenüber tauben oder schwerhörigen Personen. Audismus basiert auf „verinnerlichten Vorurteilen und der Vorstellung, dass ein Leben ohne Gehör minderwertig ist“ (Yomma 2021). Die Folgen sind diskriminierende Handlungen wie Stigmatisierung, Ausgrenzung und ungerechtfertigte Ungleichbehandlung. Viele Hörende fühlen sich gehörlosen und schwerhörigen Menschen automatisch überlegen, wodurch die Gehörlosenkultur, die Gebärdensprache und die Zugehörigkeit zu dieser ethnischen Gruppe abgewertet werden (vgl. Yomma 2021). Der Begriff „taubstumm“, der oftmals zur Beschreibung nicht-hörender Menschen verwendet wird, wird ebenfalls als diskriminierend empfunden. Gehörlose Menschen sind genauso wie Hörende in der Lage zu sprechen, entweder in der Lautsprache oder in der Gebärdensprache. Einige Begriffe, die von Menschen mit eingeschränktem Hörvermögen bevorzugt werden, sind „schwerhörig“, „hörbeeinträchtigt“, „hörbehindert“, „gehörlos“ oder auch „taub“ (um sich gegen die defizitäre Konnotation des Wortes „gehörlos“ zu wehren) (vgl. Leidmedien 2021).

Im Englischen spricht man von „deaf“, wenn die betreffende Person Sprache nicht ausschließlich auf auditivem Weg wahrnehmen kann. Die Wörter *Deaf* und *Hard of Hearing* werden dann großgeschrieben, wenn sich die gehörlose Person als Angehörige:r der Gehörlosenkultur sieht. Spätertaubte sehen sich oftmals als Teil der Hörendenkultur (vgl. Jüngst 2010: 125f.).

3.2.4 Überblick über die historische Entwicklung der SDH-Untertitelung

Die (interlinguale) Untertitelung nahm ihren Anfang bei der Übersetzung von Zwischentiteln in Stummfilmen, welche sowohl erzählerische Funktion hatten als auch Dialoge beinhalten konnten. Erst zu Ende der Stummfilmzeit konnten die Zwischentitel auch auf die Bilder übertragen werden. Die Untertitel von Tonfilmen hingegen beinhalten nur den Dialog sowie im

Bildmaterial geschriebene Informationen, zum Beispiel Titel von Zeitschriftenartikeln, also Forced Narratives (vgl. O’Sullivan & Cornu 2018: 20).

Da die ersten Tonfilme, die auf der ganzen Welt ausgestrahlt wurden, vorwiegend in den USA produziert wurden, wurden Untertitelungen vorerst hauptsächlich aus dem Englischen in andere Sprachen vorgenommen. Der Hauptmarkt für amerikanische Produktionen war Europa, weshalb die Untertitelungsprozesse zu Beginn vor allem ebendort entwickelt wurden. Zu diesen ersten Entwicklungen zählt der in den Anfangsjahren vor allem in den USA und Europa eingesetzte Fotodruck, der bei weißem Filmhintergrund aufgrund der weißen Buchstaben oft nur schwer lesbar war. Schon bald wurden Methoden zur Verbesserung der Lesbarkeit in Europa entwickelt, von denen sich vor allem chemische Verfahren sowie das Arbeiten mit Lasern bewährten. Die Untertitel wurden mithilfe dieser Methoden in den Filmstreifen eingebrannt und waren aufgrund der schwarzen Umrandungen der Buchstaben besser lesbar, wobei sie bei gänzlich weißem Hintergrund immer noch verschwinden konnten. Dieses Problem konnte erst in den 2000er-Jahren beseitigt werden, als die digitale Untertitelung perfekte Lesbarkeit auf allen möglichen Hintergründen garantieren konnte. Die elektronische Untertitelung wurde – im Gegensatz zu den vorher genannten Methoden, die hauptsächlich für Kinofilme verwendet wurden –, ab den 1970er-Jahren für das Fernsehen eingesetzt. Bei dieser Methode wurden computergenerierte Texte erstellt und in das elektronische TV-Bild eingebunden. In weiterer Folge wurde dieses Verfahren auf Filme auf VHS-Kassetten und später auf DVDs angewendet, wobei diese heute ebenfalls auf digitale Verfahren zurückgreifen (vgl. O’Sullivan & Cornu 2018: 18).

Es ist aufgrund der nur spärlich vorhandenen Dokumentation nicht klar zu sagen, wer die Untertitel-Pionier:innen sind und wie sie arbeiteten. Die frühen Talkies, deren Untertitel in und aus dem Englischen übersetzt wurden, hatten scheinbar eine relativ kleine Anzahl an Untertiteln, die wohl nur die Hauptaussage übermitteln sollten. Dies gilt vor allem für Filme, die ins Englische übersetzt wurden, und für portugiesische Untertitel für den brasilianischen Markt. Filme, die ins Französische übersetzt wurden, scheinen schon von Anfang an eine größere Anzahl an Untertiteln gehabt zu haben. Es ist denkbar, dass die Untertitelung in den 1930er-Jahren in den USA mit Herman Weinberg ihren Anfang nahm, der sich als erster Untertitler in New York bezeichnet. Der erste von ihm untertitelte Film war „Zwei Herzen im ¾ Takt“ von Géza von Bolváry (1930), wobei sich die Untertitelung erst in den darauffolgenden Jahren durchsetzte. So beschreibt Weinberg, dass er zunächst den Zuschauer:innen in Bildtiteln darlegte, was in den nächsten Minuten geschehen würde, bevor er auf die Art von Untertiteln, wie wir sie heute kennen, umstieg. Die Erfindung der Moviola Tonbearbeitungsmaschine

erlaubte es ihm, den Dialog und die Untertitel so synchron wie möglich zu schalten (vgl. O’Sullivan & Cornu 2018: 21).

Die SDH-Untertitelung wurde vor allem vom britischen Sender BBC vorangetrieben, der von Anfang an eine Vormachtstellung innehatte. Bereits im Jahr 1972 kündigte dieser den Ceefax Teletext-Service an, welcher 1979 zur Einführung von SDH-Untertiteln führte, wobei schon damals ein wöchentliches Nachrichtenprogramm mit offenen Untertiteln zur Verfügung gestellt wurde. Der flämischsprachige Teil Belgiens, Deutschland, Italien und die Niederlande folgten diesem Trend in den 1980ern; Portugal und Spanien kamen in den 1990er-Jahren dazu. Die meisten europäischen Fernsehsender entwickelten zu dieser Zeit den Teletext, der maßgeblich zur Verbreitung von SDH-Untertiteln beitrug. So wurde zum Beispiel am 1. Juni 1980 zum ersten Mal Videotext im deutschen Fernsehen gesendet, nämlich als Pilotprojekt der Medienanstalten ZDF und ARD (vgl. Remael 2007: 24).

Zu Beginn war die Anzahl der Sendungen, für die SDH-Untertitel zur Verfügung gestellt wurden, sehr klein; beim Fernsehsender ORF etwa wurden im Jahr 1980 nur fünf Stunden des wöchentlichen Programms untitled (vgl. ORF 2023). Oftmals kamen diese bei Sendungen zum Einsatz, welche von den Sendern als besonders interessant für ein hörgeschädigtes und gehörloses Publikum eingeschätzt wurden. Auch die Vorbereitungszeit war ein wichtiger Faktor, denn nicht viele Sendungen waren im Vorhinein zur Bearbeitung verfügbar. Nach und nach wurden Nachrichtensendungen mit SDH-Untertiteln ausgestattet. Der flämische öffentliche Sender VRT sendete beispielsweise 1983 seine ersten Untertitel im Rahmen von Sportsendungen, und 1992 wurden die Untertitel im Zuge von zwei Nachrichtensendungen zu einem festen Bestandteil des Senders (vgl. Remael 2007: 24).

Trotz der Veränderungen im Hinblick auf die Barrierefreiheit der Inhalte waren die Bedürfnisse der Hörgeschädigten- und Gehörlosencommunitys noch nicht erfüllt, denn es waren vor allem die öffentlichen Sender, die SDH-Untertitelungen erstellten, während sich die privaten Sender eher zurückhielten. Die einzige Ausnahme war das Vereinigte Königreich, wo sowohl private als auch öffentliche Sender Untertitel zur Verfügung stellten. Die Vergrößerung des barrierefreien Angebots ging meistens mit einer intensiveren Beschäftigung und Zusammenarbeit mit der Zielgruppe sowie mit der Entwicklung neuer Untertitelungstechnologien und -arten, etwa der Live-Untertitelung, einher. Die Schnelligkeit, mit der diese Entwicklungen sich vollzogen, war zumeist mit der Einführung neuer Vorschriften oder anderer Vereinbarungen zwischen Regierungen und Medienanstalten verbunden. Diese rechtlichen Veränderungen wurden wiederum vor allem von Gehörlosenorganisationen vorangetrieben, welche vermehrt Druck ausübten, um die

Dringlichkeit der Zugänglichkeit der Inhalte deutlich zu machen. Europaweite Initiativen, zum Beispiel die Richtlinie „Fernsehen ohne Grenzen“, seien laut diverser Forscher:innen zwar wichtig gewesen, aber bei Weitem nicht ausreichend für den Ausbau des barrierefreien Angebots (vgl. Remael 2007: 25).

Heute wird die SDH-Untertitelung als eines der wertvollsten Mittel für den barrierefreien Zugang zu Information und Unterhaltung für hörbehinderte Menschen angesehen. Ein ausschlaggebender Punkt für diese Entwicklung war das am 3. Mai 2008 beschlossene Übereinkommen der Vereinten Nationen über die Rechte von Menschen mit Behinderungen (UN-Behindertenrechtskonvention, s. Kapitel 2.1.1). In Artikel 9 der Konvention liegt das Hauptaugenmerk darauf, Menschen mit Behinderung den Zugang zu Information und Kommunikation zu ermöglichen, um ihnen die Möglichkeit eines selbstbestimmten Lebens und der Teilnahme an allen Aspekten des Lebens zu geben. Diese soziopolitische Entwicklung führte zu der Verbreitung neuer Services für behinderte Menschen auf nationalem Level, etwa die Durchsetzung neuer Gesetze und Verordnungen – auch in Bezug auf den Zugang zu Information, vor allem durch Fernsehen und Internet. Diese Rahmenbedingungen hatten ebenfalls direkten Einfluss auf barrierefreie Services, darunter Untertitel, Audiodeskription und Gebärdensprachdolmetschen, und führte zur Einführung von SDH-Untertitelung in vielen Ländern bzw. zur weiteren Verbreitung dieser (vgl. Neves 2018: 84).

Barrierefreie Services haben sich vor allem dort weitgehend etabliert, wo die wichtigsten Stakeholder (Gesetzgeber:innen, Anbieter:innen, Produzent:innen, Vertreiber:innen und Endnutzer:innen) durch gemeinsame Anstrengungen Veränderungen bewirken konnten. Im Fall der SDH-Untertitelung ist es vor allem die Gehörlosencommunity, die die Verbesserung von Qualität und Quantität der Services hervorgerufen hat. Die Communitys haben sich mehr und mehr für ihre Ziele eingesetzt und wurden schließlich zu einer treibenden Kraft, welche zum Beispiel Untertitelungsstandards festlegte. So zum Beispiel auch in Spanien, wo verschiedene Organisationen gemeinsam den UNE 153010-Standard entwickelten. Gehörlosenverbände haben außerdem in nationalen und internationalen Forschungsprojekten wie SAVAS (Sharing AudioVisual Language Resources for Automatic Subtitling), SUMAT (Online Services for Subtitling by Machine Translation) oder den bereits genannten Projekten DTV4ALL und HBBTV4ALL mitgewirkt. Diese Lobbyarbeit kann sowohl durch die Teilnahme an diesen Studien als auch durch individuelle Arbeit, zum Beispiel in sozialen Netzwerken, durchgeführt werden. Außerdem haben sich Gehörlosenverbände aktiv gegen die Diskriminierung von Netflix, Fox, Universal, Warner Bros. und Paramount

eingesetzt, da diese ihrer Meinung nach nicht genug SDH-Untertitel zur Verfügung gestellt hatten. Derartige Aktionen, die im Jahr 2012 mit Petitionen und Klagen in Bezug auf die englische Sprache angingen, haben sich bis heute auch auf weniger verbreitete Sprachen wie Griechisch oder Portugiesisch übertragen. Diese Entwicklungen zeigen, dass Gehörlosencommunitys sich im Laufe der Zeit immer mehr Bewusstsein über ihre Rechte aneigneten und diese auch durchsetzten (vgl. Neves 2018: 85f.).

Auch die wichtige Rolle der Endnutzer:innen in der Untertitel-Versorgungskette wurde durch diese Entwicklungen unterstrichen. Gehörlose Zuseher:innen nehmen aktiv an der Verbesserung von Qualität und Quantität teil, indem sie sich an verschiedenen Levels innerhalb dieser Versorgungskette beteiligen. Dennoch muss erwähnt werden, dass kein Service für alle Endnutzer:innen passend sein kann, da die Teilnehmer:innen nur einen Teil der Menschen darstellen, die das Untertitelungsangebot aus Gründen der Barrierefreiheit benutzen. Außerdem kann die Adäquatheit der Untertitel nur mithilfe von objektiven Daten, die etwa durch Eyetracker, Magnetresonanztomographien oder Elektroenzephalogramme gesammelt worden sind, beurteilt werden, nicht aber durch subjektive Einschätzungen (vgl. Neves 2018: 86). Neves (2018: 86) unterstreicht außerdem in Anlehnung an Romero-Fresco (2015): „What viewers think of SDH, how they understand these subtitles and how they view them‘ does not necessarily match.“

Neves (2018) identifiziert darüber hinaus in der heutigen Zeit drei verschiedene Schnelligkeiten der Bereitstellung von SDH-Untertiteln für audiovisuelle Inhalte, die vor allem für das Fernsehen gelten, da dieses die bevorzugte Plattform für barrierefreien Zugang ist. Die Pionierstaaten, wie die USA, Kanada, das Vereinigte Königreich, Australien und Dänemark, befinden sich im Moment bei fast 100 Prozent des audiovisuellen Angebots mit SDH-Untertiteln, welche aus einem Mix zwischen vorgefertigten und Live-Untertiteln bestehen. Medienhäuser wie die BBC, die im Fernsehen bereits 100 Prozent erreicht haben, arbeiten daran, auch ihre Mediatheken vollständig barrierefrei zu gestalten. Sobald auch dieser Schritt erreicht ist, sollen die bereits erstellten Untertitel auch auf andere Medienformate übertragen werden. Diese Dynamik ist laut Neves (2018) typisch für audiovisuelle Inhalte, denn sie folgt dem Schema „make available“ – „increase quantity“ – „address quality“. Im darauffolgenden Schritt soll die Diversität eine Hauptrolle spielen, wobei auch hier die Aspekte Quantität und Qualität beachtet werden sollen (vgl. Neves 2018: 84f.).

Die zweite Kategorie an Ländern, zu denen zum Beispiel Italien, Spanien, Portugal und Deutschland gehören, hat ihr Angebot an barrierefreien Inhalten ebenfalls – sowohl auf privaten als auch auf öffentlichen Plattformen – erweitert. Diese seit den 1980er- und 1990er-Jahren

vorhandenen Anstrengungen sind in diesen Ländern immer dann erfolgreicher, wenn Lobby- oder Forschungsarbeit betrieben wird, die die Entwicklungen vorantreibt. In diesen Ländern kamen auch neue, experimentelle Zugänge zur SDH-Untertitelung auf – zum Beispiel in Portugal, wo der Fernsehsender TRP gänzlich auf Respeaking verzichtet und Live-Untertitelung durch automatische Spracherkennung anbietet. In Spanien wiederum wurde der bereits genannte UNE 153010-Standard entwickelt, welcher von den bisherigen Traditionen abweicht und neue Konventionen in Bezug auf Farbgebung und Platzierung von Untertiteln schafft (vgl. Neves 2018: 85).

Die dritte Gruppe an Ländern bzw. Regionen ist jene, in denen die SDH-Untertitelung erst jetzt, erneut angetrieben durch Forschung und Veränderungen in der Industrie, langsam beginnt sich durchzusetzen. Zu diesen Regionen zählen Südamerika, Asien, der Süden Afrikas und der Nahe Osten (vgl. Neves 2018: 85).

3.2.5 SDH-Richtlinien

Im Moment gibt es noch keine vereinheitlichten, international gültigen Standards für SDH-Untertitelungen, weshalb es oftmals zu Unklarheiten auf Seiten der Rezipient:innen und auch zu Inkonsistenzen bei den Untertitelungen kommen kann. Aus diesem Grund fordern Expert:innen schon lange eine Vereinheitlichung der Untertitelungsregelungen bis zu einem gewissen, realisierbaren Punkt, diese kam jedoch bisher (unter anderem aufgrund von fehlendem Konsens zwischen den Medienanstalten und der variierenden technischen Entwicklungen) noch nicht zustande. Harmonisierte Standards, die einen funktionellen oder technischen Aspekt reglementieren (zum Beispiel eingesetzte Symbole in den Untertiteln), seien laut Neves (2005) leichter zu erreichen als die Untertitelnormen an sich, da die Rezipient:innen bereits die in ihrem (Sprach-)Gebiet vorherrschende Art und Weise der Untertitelung gewohnt sind. Es sei sehr schwer, diese Gewohnheiten zu ändern, und außerdem hängen diese auch von Faktoren wie nationale Identität sowie Bildungs- und sozialem Hintergrund ab. Die Vereinheitlichung solle sich also, so Expert:innen, in einem praktischen Rahmen bewegen und darüber hinaus idealerweise von einer Aufsichtsbehörde zur Qualitätssicherung geprüft werden (vgl. Remael 2007: 39-40).

Einer der Unterschiede in Bezug auf Untertitelungsrichtlinien im deutsch-, englisch- und französischsprachigen Raum ist die Verwendung von Farben. Im deutsch- und englischsprachigen Raum werden Schriftfarben in den meisten Fällen als Sprecher:innenidentifikationen eingesetzt. Hier ist jedoch zu notieren, dass im

deutschsprachigen Raum die Farben Weiß, Gelb, Grün, Cyan und Magenta zur Verfügung stehen (vgl. Das Erste 2021), während im englischsprachigen Raum Magenta nicht verwendet wird (vgl. BBC 2021). Im Gegensatz dazu wird im französischsprachigen „Dossier sous-titrage“ des Conseil Supérieur de l’Audiovisuel (CSA) folgender Einsatz von Farben vorgeschrieben:

- Weiß bei einer am Bildschirm sichtbaren sprechenden Person
- Gelb bei einer nicht am Bildschirm sichtbaren (in der Szene vorkommenden) sprechenden Person
- Rot bei Geräuschangaben
- Magenta für die Angabe von Musik
- Cyan für inneren Monolog oder Off-Sprecher:innen
- Grün für die Verwendung von Fremdsprachen (vgl. CSA 2012: 12).

Im französischen Sprachraum stehen folglich die meisten Farben zur Verfügung, da hier auch Rot verwendet wird. Die Farbgebung gestaltet sich also deutlich anders als im deutsch- und englischsprachigen Raum. Diese Unterschiede könnten bei Rezipient:innen ohne Kenntnis der jeweiligen Untertitelungstraditionen und -konventionen möglicherweise Verwirrung auslösen.

Dieses Beispiel soll illustrieren, dass der Vergleich von im Fernsehen ausgestrahlten Programmen sich eventuell schwierig gestalten würde, da schon in Bezug auf diesen einen Aspekt ein so deutlicher Unterschied zwischen den Sprachräumen besteht. Im Gegensatz dazu sind die Netflix-Richtlinien einheitlich für alle Sprachen und Zielgruppen (intra- und interlingual) gestaltet und decken sich größtenteils mit den gängigen Untertitelungskonventionen, wie sie zum Beispiel bei Díaz Cintas und Remael (2021) genannt werden. In den nächsten Absätzen sollen die generellen Aspekte der Netflix-Untertitelungsrichtlinien genauer dargelegt werden.

Die Untertitel müssen demnach eine Mindestdauer von fünf Sechsteln einer Sekunde haben – das bedeutet, dass zum Beispiel bei einer Framerate von 24 Frames pro Sekunde ein Untertitel mindestens 20 Frames lang sein muss. Darüber hinaus gilt eine Maximaldauer von sieben Sekunden pro Untertitel. Ein einzelner Untertitel darf des Weiteren aus maximal zwei Zeilen bestehen, wobei diese Aufteilung nur zulässig ist, wenn die maximale Zeichenanzahl pro Zeile überschritten wird. Für die Aufteilung in zwei Untertitel werden einige Grundprinzipien vorgegeben: So soll der Zeilenumbruch nach Satzzeichen, aber vor Konjunktionen oder Präpositionen erfolgen. Im Umkehrschluss soll ein Zeilenumbruch aber keinesfalls Sinneinheiten (z.B. Nomen und Artikel) trennen (vgl. Netflix 2021a).

Die Untertitel sollen in allen Sprachen zentriert sein und sich entweder am unteren oder am oberen Bildrand befinden; nur im Japanischen ist das Verwenden vertikaler Untertitel erlaubt. Die anderssprachigen Untertitel sind nur dann aus ihrer Ausgangsposition am unteren Bildrand nach oben zu verschieben, wenn sie ansonsten handlungsrelevante Elemente am unteren Bildrand verdecken würden. Dazu zählen etwa Forced Narratives oder Einblendungen. Wenn auf beiden Seiten schon Text vorhanden ist, müsse entschieden werden, wo der Untertitel besser zu lesen ist (vgl. Netflix 2021a).

Zusätzlich zu diesen allgemein gehaltenen Vorgaben und Empfehlungen existieren Regelungen für die jeweiligen Sprachen. Es sei hier angemerkt, dass sich ein Großteil der Vorgaben für die deutschen und französischen Sprachfassungen auf interlinguale Untertitelungen bezieht und nur ein kleiner Teil den SDH-Untertiteln gewidmet wird, während im Englischen eine ausführlichere Darstellung der Regeln angegeben wird. Im folgenden Teil werden hauptsächlich die Vorgaben für die englische Sprachfassung („English Timed Text Style Guide“) behandelt, da diese sich größtenteils mit den beiden anderen Guides decken. Auf etwaige Abweichungen werde ich explizit hinweisen.

Die maximale Zeichenanzahl pro Zeile beträgt 42 Zeichen. Ein Untertitel bis zu dieser Zeichenanzahl muss demnach ein Einzeiler bleiben, während er ab 43 Zeichen auf zwei Zeilen aufgeteilt werden muss. Für den Fall, dass zwei Zeilen verwendet werden, soll die untere Zeile, wenn möglich, länger sein als die obere, wobei die obere Zeile wiederum nicht aus weniger als drei Wörtern bestehen soll. Die Lesegeschwindigkeit für Programme mit erwachsener Zielgruppe beträgt 20 Zeichen pro Sekunde, während bei Kindersendungen nur 17 Zeichen pro Sekunde zulässig sind. Diese Angaben dürfen im Normalfall um maximal 20 Prozent überschritten werden (vgl. Netflix 2021b).

Generell gilt, dass so viel wie möglich von dem Originaldialog im Untertitel wiedergegeben werden muss und auch möglichst keine Vereinfachung stattfinden soll. Wenn Inhalte in die jeweilige Sprache übersetzt und synchronisiert worden sind, sollen die Untertitler:innen sich möglichst auf die synchronisierte Sprachfassung berufen. Das Kürzen des Originaldialogs ist nur dann zulässig, wenn die Lesegeschwindigkeit oder die Synchronizität mit dem Ton ein Problem darstellen. Wenn ein Satz aufgrund der Lesegeschwindigkeit verändert werden soll, soll zuerst versucht werden, den Text zu reduzieren, zu löschen oder zu verdichten, bevor paraphrasiert wird. Die Transkription des Dialogs soll außerdem möglichst wortgetreu wiedergegeben werden, auch in Bezug auf die Satzstruktur; Dialekte und Slangausdrücke sollen beibehalten werden. Außerdem soll der Originaldialog niemals zensiert werden. Wenn das Adjektiv „black“ für die Bezeichnung einer Person, Personengruppe oder

Institution verwendet wird, soll dieses groß geschrieben werden; Diese Vorgabe gilt auch für die deutsche Sprache (vgl. Netflix 2021c), jedoch nicht im Französischen (vgl. Netflix 2021d). Im Englischen sollen auch „deaf“ und „indigenous“ in diesem Fall großgeschrieben werden. Absichtliche Fehler im Originaldialog sollen nur dann wiedergegeben werden, wenn sie handlungsrelevant sind. In allen drei Sprachen sollen Zahlen von eins bis zehn ausgeschrieben, alle Zahlen darüber in Ziffern geschrieben werden (vgl. Netflix 2021b).

Der Sprecher:innenwechsel soll immer mithilfe eines Bindestrichs ohne Leerzeichen erfolgen. Wenn zusätzlich eine Speaker-ID notwendig ist, steht diese zwischen dem Bindestrich und der Aussage. Bindestriche können auch verwendet werden, um etwa eine Aussage und ein Geräusch bzw. zwei Geräusche voneinander zu trennen, die nicht von derselben Person/Quelle kommen. Wenn die Aussage und das Geräusch bzw. beide Geräusche von derselben Quelle stammen, ist kein Bindestrich nötig. Auch die Kursivsetzung ist in allen Sprachen gleich: So sollen etwa Off-Sprecher:innen, Gedanken und innere Monologe kursiv gesetzt werden. Auch fremdsprachige Wörter, Liedtexte und jegliche technikgestützte Kommunikation, etwa Telefonate, sollen kursiv gesetzt werden. Es ist außerdem möglich, durch Kursivsetzung eine Hervorhebung durchzuführen (vgl. Netflix 2021b).

Speaker-IDs oder Soundeffekte sollen durch eckige Klammern dargestellt werden. Hierbei ist darauf zu achten, dass – mit Ausnahme von Nomen im Deutschen – alle Wörter klein geschrieben werden. Diese Angaben sollen nur erfolgen, wenn sie nicht auf dem Bildschirm erkennbar sind. Außerdem werden auch bei vollen Sätzen keine Satzzeichen gesetzt. Bei Speaker-IDs für bisher noch nicht namentlich bekannte Personen ist darauf zu achten, dass die Personen auch nicht in den Untertiteln identifiziert werden. In solchen Fällen sollen Bezeichnungen wie [Mann], [Frau] oder [Ärztin] verwendet werden. Nicht handlungsrelevante Hintergrundmusik soll durch generische Musik-Untertitel, etwa [rock music playing over stereo], beschrieben werden, während bei handlungsrelevanter Musik auch Titel und Interpret genannt werden, zum Beispiel [Musik: "Forever Your Girl" von Paula Abdul]. Songtexte sollen mithilfe einer Musiknote (♪) mit Leerzeichen davor bzw. danach am Anfang und am Ende jedes Untertitels gekennzeichnet werden. Auch handlungsrelevante Soundeffekte sollen in eckigen Klammern erwähnt werden, wenn sie nicht sichtbar sind – hierzu zählt auch Stille, wenn zum Beispiel ein Lied abrupt endet. Die Soundeffekte sollen außerdem so detailliert wie möglich durch Adjektive und Adverbien beschrieben werden, wenn es zeitlich und räumlich möglich ist. Bei fremdsprachigem Dialog gibt es zwei Möglichkeiten: Wurde der Dialog übersetzt, sollen Formulierungen wie [Spanisch] oder [in Spanisch] auf die Sprache hinweisen. Wenn der Dialog nicht verstanden werden soll, werden Formulierungen

wie [spricht Spanisch] benutzt. Die gesprochene Sprache soll jedoch immer identifiziert werden; Formulierungen wie [spricht Fremdsprache] sind nicht zulässig (vgl. Netflix 2021c).

3.3 Aktuelle Forschungsschwerpunkte

3.3.1 Lesbarkeit, Lesegeschwindigkeit und Verständlichkeit

Auch in der heutigen Forschung gibt es einige Diskussionsschwerpunkte, die sich mit SDH-Untertitelung befassen. Einer von ihnen ist die ewig diskutierte Frage, ob nun wörtliche (*verbatim subtitles*) oder bearbeitete (*edited subtitles*) Untertitel bevorzugt werden sollen. Einerseits präferieren ertaubte und hörgeschädigte Rezipient:innen, welche auf ihr restliches Hörvermögen und auf Lippenlesen zurückgreifen, wörtliche Untertitel. Dieser Zugang wird auch von Gehörlosenverbänden und Sendeanstalten unterstützt. Andererseits sprechen sich Forscher:innen immer wieder für bearbeitete Untertitel aus. Es wird jedoch betont, dass die Lesekompetenz der einzelnen Zuseher:innen ausschlaggebend für die Effektivität der Untertiteltypen ist. Es sei daher umso wichtiger, dass zuerst festgestellt wird, wie gehörlose und hörbehinderte Menschen generell lesen, bevor verstanden werden kann, wie Untertitel verfasst und präsentiert werden sollen, um deren Verständnis zu fördern. In diesem Zusammenhang soll auch erwähnt werden, dass die Lesekompetenz auch nach Art des Textes variieren kann, da die Art zu lesen ebenfalls von dem Medium abhängt, von dem man abliest. Untertitel können schwerer zu lesen sein als statischer Text auf Papier oder auch auf Bildschirmen, da sie vom bewegten Bild beeinflusst werden können. Außerdem haben Leser:innen möglicherweise wenig bis keine Zeit, die Untertitel erneut zu lesen oder zu recherchieren. Auch aus diesem Grund sind Studien zur Lesbarkeit wichtig, um diese so gut wie möglich an die Bedürfnisse der Community anzupassen (vgl. Neves 2018: 89f.).

Nicht nur die Lesbarkeit, sondern auch die Rezeption von SDH-Untertiteln ist bis heute einer der großen Forschungsschwerpunkte. Die größte SDH-Rezeptionsstudie bisher war das von der EU geförderte Projekt DTV4ALL, das die neuen technischen Möglichkeiten des digitalen Fernsehens nutzen wollte, um hör- und sehbehinderten Menschen den Zugang zum Fernsehen zu erleichtern (vgl. Romero-Fresco 2015). Das Projekt dauerte von 2008 bis 2011 und wurde in Dänemark, Polen, dem Vereinigten Königreich, Spanien, Italien, Frankreich und Deutschland durchgeführt. Neben Erkenntnissen auf nationaler Ebene wurden auch nationenübergreifende Daten gesammelt, die zeigten, wie verschiedene Arten von Rezipient:innen SDH in verschiedenen Ländern und Sprachen verarbeiten und verstehen. Es wurde festgestellt, dass, obwohl gehörlose Zuseher:innen die Untertitel um zehn Prozent

schneller auf dem Bildschirm erkennen und sie um acht Prozent länger lesen als hörende Zuseher:innen, ihr Verständnis der Untertitel um 15 Prozent schlechter ist, was ein Indiz für Leseschwächen sein kann. Und obwohl sie weniger Zeit mit dem Betrachten des Bildmaterials verbringen, ist das visuelle Verständnis der gehörlosen Zuseher:innen genauso gut bzw. teilweise sogar höher als das der Hörenden. Gehörlose scheinen also ihre teils schlechteren Lesefähigkeiten mit gutem Sehvermögen und visuellem Verständnis auszugleichen (s. hierzu das nächste Unterkapitel). Das Projekt erzielte außerdem Erkenntnisse in Bezug auf Lesegeschwindigkeiten, welche mithilfe von Eyetracking erfasst wurden. Die Lesegeschwindigkeit ist jene Geschwindigkeit, die ein:e bestimmte:r Zuseher:in benötigt, um audiovisuelles Material anzusehen (inklusive Untertitel, Bildmaterial und, wenn möglich, Tonmaterial). Während die traditionelle Empfehlung von 150 Wörtern pro Minute (wpm) eine gleichmäßige Aufmerksamkeitsverteilung zwischen Untertiteln und Bildmaterial ermöglicht, erhalten die Untertitel bei 180 wpm mehr Aufmerksamkeit als das Bildmaterial (vgl. Romero-Fresco 2019b: 551).

3.3.2 Informationsdichte und Redundanz

Multimodalität wird für verschiedenste Ausprägungen von Translation als besondere Herausforderung gesehen: Der „bloße“ Sprachtransfer wird durch das Zusammenspiel verschiedener Modi ausgeweitet und dadurch erschwert. Doch O’Sullivan (2013) betont, dass Multimodalität auch eine wichtige Ressource sein kann – sowohl für Translator:innen als auch für Rezipient:innen dieser Texte. Zum Beispiel können Bilder in Terminologiedatenbanken Übersetzer:innen dabei helfen, (fachsprachliche) Terminologie zu verstehen und die richtige zielsprachige Entsprechung zu finden. Im Fall der Untertitelung ist das Zusammenspiel von Ton, Bild und Text ebenfalls sowohl Herausforderung als auch Hilfestellung: Einerseits können Beschränkungen, wie Platzmangel und festgelegte Zeichenanzahl, den Untertitelungsprozess erschweren. Andererseits wiederum kann genau dieses Zusammenspiel dabei helfen, in bestimmten Situationen genug Kontext zu liefern, sodass verbale Elemente bis zu einem bestimmten Punkt überflüssig werden, wodurch das Verdichten des Textes einfacher gemacht wird (vgl. O’Sullivan 2013: 11).

Diese Auffassung teilen auch Tiittula & Hirvonen (2019: 16): „Wegen der Multimodalität des Gesamten ist es [...] nicht unbedingt notwendig, jedes nicht wahrnehmbare Element zu übersetzen oder zu ersetzen, wenn die Bedeutung durch andere Modalitäten erschließbar ist.“ Narrative Informationen werden in audiovisuellen Produkten durch die

Verbindung unterschiedlicher Modalitäten übermittelt. Wenn einer der dafür benötigten Sinneskanäle nicht zugänglich ist und somit eine Modalität nicht verwendet werden kann, können andere Modalitäten an Relevanz gewinnen. Tiittula & Hirvonen (2019) führen hier etwa die Rolle der Musik bei Hörfilmen an, welche in einem audiovisuellen Produkt eventuell weniger Aufmerksamkeit bekommen würde, da sie sich diese mit den Bildinformationen teilen müsste. Die intermodale Übersetzung, zu der auch die SDH-Untertitelung gehört, zeichnet sich dadurch aus, dass eine Modalität durch eine andere ersetzt wird, wobei einige Elemente einfacher zu ersetzen sind als andere. So ist es im Vergleich schwieriger, zum Beispiel Geräusche und Musik oder auch Sprachvarietäten durch geschriebene Sprache zu reproduzieren (vgl. Tiittula & Hirvonen 2019: 15f.).

Wie bereits diskutiert, ist für Hörgeschädigte und Gehörlose einer der bedeutungstragenden Modi audiovisueller Texte nicht zugänglich. Das muss jedoch, so Remael und Reviers (2018), nicht zwingend ein Problem darstellen: Die Rezipient:innen sind sich dessen bewusst und wenden aktiv Kompensationsstrategien an, indem sie auf externe Quellen, etwa ihr eigenes Weltwissen, zurückgreifen. Eine von Remael und Reviers (2018) durchgeführte Studie zeigte, dass weniger Redundanz einerseits zu einem größeren Bedürfnis nach Rekonstruktion und zu höherer mentaler Anstrengung beim Publikum führe. Andererseits könne sie auch eine Hilfestellung für hörgeschädigte Personen sein, da diese in informationsreichen Szenen auf konzise Untertitel zurückgreifen können und so das Risiko für Informationsüberlastung senken können (vgl. Remael & Reviers 2018: 277).

Auch Vandaele (2018: 234) betont, dass Kürzungen nicht immer einen negativen Effekt auf die Untertitelung haben müssen: Die Prägnanz kann es den Zuschauern ermöglichen, mehr auf die non- und paraverbalen Merkmale der Aussage zu achten (vgl. Vandaele 2018: 234.).

3.3.2 Integration von Symbolen

Auch Verständnishilfen in Form von Symbolen oder Gebärdensprache in Untertiteln werden seit einiger Zeit erforscht. Dabei gibt es verschiedene Möglichkeiten, wie Symbole in Untertitel eingebaut werden können: ein Beispiel hierfür wären Icons für Geräusch- und Tonkontexte, etwa ein klingelndes Telefon oder eine Lautsprecherdurchsage (vgl. Civera 2005). Eine in Portugal durchgeführte Studie zeigte darüber hinaus, dass viele Rezipient:innen dem Verwenden von Smileys zum Anzeigen von Emotionen positiv gegenüberstanden (vgl. Neves 2005).

Schwieriger wird es bei Sprecher:innen- und Emotionsdarstellungen. Damit keine echten Gesichter verwendet werden müssen, weil dies zum Beispiel aus Datenschutzgründen problematisch sein kann, schlagen Forschende die Verwendung sogenannter Face Alive Icons (FAIs) vor. FAIs sind Fotos von echten Personen, die mithilfe einer Software verändert werden. Das Erkennen der dargestellten Emotionen ist aber ebenfalls eine Herausforderung, da in Studien nicht jede Emotion von den Proband:innen einfach erkannt werden konnte. Die Emotion Ekel zum Beispiel wurde nur von einem Bruchteil der Personen richtig identifiziert (vgl. Civera & Orero 2010: 159f.). In anderen Untersuchungen betrachteten Zuseher:innen diesen Zugang mit Skepsis, da sie es nicht für möglich hielten, eine große Bandbreite an Emotionen durch einfach erkennbare Smileys darzustellen (vgl. Romero-Fresco 2015), was die zuvor genannten Ergebnisse von Civera & Orero (2010) unterstreicht.

3.3.3 Automatisierungsprozesse

Die Sprachindustrie ist außerdem immer auf der Suche nach neuen, wirtschaftlicheren Optionen, um die barrierefreien Services günstiger und damit leichter zugänglich zu machen, was mehrere große, von der EU geförderte Projekte zur Folge hatte. Eines davon war SUMAT (SUBtitling by MACHine Translation), welches sich auf interlinguale Untertitelung konzentrierte und die Möglichkeit von maschineller Übersetzung und Postedition für ebendiese untersuchte. SAVAS (Sharing AudioVisual language resources for Automatic Subtitling) wiederum hatte das Ziel, automatisiert intralinguale Untertitel durch Spracherkennung herzustellen. Projekte dieser Art funktionieren immer durch den Zusammenschluss verschiedener Gruppen, darunter Unternehmen aus der Sprachindustrie, welche die immensen Datenmengen zur Verfügung stellen können, die für jegliche maschinell unterstützte Übersetzung nötig sind, Softwareentwickler:innen und Universitätsabteilungen mit Schwerpunkt auf Computerlinguistik und ähnlichen Feldern. Ein Vorteil von solchen groß angelegten Projekten ist, dass eine Vielzahl an Expert:innen mithilfe einer großen Datenmenge an neuen Ideen arbeiten kann (vgl. Remael et al. 2018: 68).

Die Automatisierung von Untertitelungs- und Audiodeskriptionsprozessen (AD) ist auch bei Kurch (2018) ein wichtiges Thema. Er schlägt gewisse Teilautomatisierungsprozesse mithilfe von Sprachtechnologien vor, die sowohl für SDH als auch für AD in der Postproduktion eingesetzt werden können. In diesem Szenario wird ein maschinelles Lernsystem mit Daten bereits erstellter Untertitelungen und Audiodeskriptionen trainiert, sodass das System die gewonnenen Daten auf neues audiovisuelles Material anwenden kann.

Damit können die neuen gesprochenen Inhalte sowohl transkribiert als auch zeitlich zugeordnet werden. Das bedeutet, dass Transkription und Zeitkodierung einer Audiodatei inklusive Metadaten automatisch als SRT-Datei generiert werden können und bereits Informationen zu Interpunktion, Groß-/Kleinschreibung, Sprecher:innen-IDs und anderen Informationen enthalten sein können, was wiederum eine immense Zeit- und Arbeitsersparnis für Ersteller:innen von SDH-Untertiteln und Audiodeskriptionen zur Folge hätte (vgl. Kurch 2018: 447f).

Aus dieser SRT-Datei können dann die jeweiligen Einheiten extrahiert werden, die für die Erstellung von SDH bzw. AD benötigt werden. Auch hier ist die maschinelle Vorarbeit oft nicht fehlerfrei und muss kontrolliert werden, außerdem müssen Zusatzinformationen wie Musik und Geräusche extra in die SRT-Datei eingefügt werden. Diese Teilautomatisierung des Untertitelungsprozesses würde aber auf alle Fälle eine große Zeit- und Arbeitsersparnis für SDH-Untertitelungen bedeuten (vgl. Kurch 2018: 449ff.).

3.3.4 Kinder als Zielgruppe

Die SDH-Untertitelung für Kinder wurde unter anderem von Lorenzo (2010) anhand einer Rezeptionsstudie untersucht und zeigte, dass viele der prä- und postlingual ertaubten Proband:innen Untertitel als sehr hilfreich für ihr Verständnis empfanden und auch die Verwendung von Farbe(n) zur Sprecher:innenidentifikation hilfreich war. Mehr als die Hälfte von ihnen hätte außerdem nach eigener Aussage mehr Zeit gebraucht, um die Untertitel zu lesen – im Gegensatz zu nur 21 Prozent der hörenden Kontrollgruppe. Interessant ist, dass die Hälfte der gehörlosen Proband:innen die Vermittlung von Emotionen via Emoticons präferierte. Außerdem empfanden diese die Beschreibung von Hintergrundgeräuschen als hilfreich, wobei Onomatopoesie (z.B. „wuff“) gegenüber Beschreibungen („Bellen“) bevorzugt wurde. Einer der größten Unterschiede zwischen den hörenden und gehörlosen Kindern bestand in der Verständlichkeit des Vokabulars: Während 98 Prozent der Hörenden das Vokabular als einfach beschrieben und keine weitere Erklärung dazu brauchten, empfanden nur zwei Drittel (69 bzw. 68 Prozent) der prälokutiv bzw. postlokutiv Ertaubten das Vokabular als einfach; 55 Prozent von ihnen brauchten weitere Erklärungen (vgl. Lorenzo 2010: 133ff.).

Wünsche (2022) behandelte in ihrer Dissertation ebenfalls SDH für Kinder und führte eine groß angelegte Rezeptionsstudie mit rund 200 Teilnehmenden in Deutschland durch. Diese Studie kam teilweise zu anderen Ergebnissen als die zuvor genannte: So wurde schwierige bzw. unbekannte Lexik als einer der Hauptgründe für schlechte Verständlichkeit von Untertiteln

genannt, jedoch gaben nur 18 Prozent der Befragten an, Verständnisprobleme gehabt zu haben. Es ist demnach nicht unbedingt nötig, Vokabular, das nicht direkt dem alltäglichen Umfeld der Kinder entspricht, zu vereinfachen, wie es von einigen Richtlinien (z.B. KiKA) verlangt wird. Ein interessanter Aspekt war auch die Farbvergabe zur Sprecher:innenidentifikation, denn während die meisten Teilnehmenden sich für farbliche Untertitel aussprachen, führten in weiß gehaltene Untertitel nicht zu einem schlechteren Verständnis. Wichtiger sei, so Wünsche (2022), die Zuordenbarkeit anhand des Ausgangskommunikats (z.B. gute Sichtbarkeit der sprechenden Person). Es stellte sich außerdem heraus, dass die befragten Kinder eher eine höhere Zeichenanzahl pro Sekunde gegenüber einer niedrigeren präferierten. So wurden 12 CPS als hilfreicher empfunden als 9 CPS, was laut Wünsche (2022) daran liegen könnte, dass bei einer geringeren Zeichenanzahl auch mehr Kürzungen vorgenommen werden müssen und die Untertitel somit stärker von der Tonspur abweichen⁵. Dies könnte zu einem *split attention effect* führen, welcher die kognitive Anstrengung erhöhen und so zu schlechterem Verständnis führen könnte. Vor allem ältere Kinder, welche sich in der Untersuchung auch als Hauptnutzer:innen von SDH herausstellten, profitierten von einer höheren Zeichenanzahl (vgl. Wünsche 2022: 360ff.).

3.3.5 Sprachenlerner:innen als Zielgruppe

Auch Sprachenlerner:innen wurden als Zielgruppe für Untertitel identifiziert – so können sowohl das Konsumieren audiovisueller Produkte in einer Zweitsprache (L2) mit Untertiteln als auch das eigene Erstellen zum besseren Verständnis und zur Erweiterung des Vokabulars in der L2 führen. Dabei wurden verschiedene Arten von Untertiteln untersucht: „klassische“ Untertitel (Audio in L2, Untertitel in L1), bimodal (Audio und Untertitel in L2), und „umgekehrte“ (Audio in L1, Untertitel in L2) (vgl. Incalcaterra McLoughlin 2018:483).

Eine der ersten Studien zu intralingualen Untertiteln für Sprachenlerner:innen wurde bereits 1988 durchgeführt und berichtete von einer eher negativen Einstellung gegenüber dem Einsatz von untertitelten Programmen, da sie als Ablenkung und als Behinderung für den Lernprozess gesehen wurden. Die Untersuchung führte jedoch zu dem Ergebnis, dass alle Proband:innen schnelles Sprechen sowie unbekanntes Vokabular besser verstanden und auch viel von dem Erlernten behalten wurde (vgl. Vanderplank 1988). Und auch interlinguale Untertitel haben sich als hilfreich für Sprachenlerner:innen aus diversen Altersgruppen und mit

⁵ Diese Zeichenanzahl ist im mittleren Bereich der Lesegeschwindigkeiten anzusiedeln – es wäre daher durchaus interessant, wie sich eine höhere Zeichenanzahl in diesem Kontext auf das Verständnis auswirken würde, wie sie etwa bei Netflix verwendet wird.

variierender Sprachkompetenz erwiesen. So zeigte zum Beispiel Danan (1992), dass der beste Abruf von gelernten Informationen mit „umgekehrten“ Untertiteln möglich war, während bimodale Untertitel gegenüber klassischen Untertiteln bevorzugt wurden. Eine große Studie mit über 2000 Teilnehmenden in Italien zeigte, dass bimodale Untertitel vor allem für Personen mit mittlerer bis hoher Sprachkompetenz geeignet waren; die meisten Befragten bevorzugten jedoch interlinguale Kombinationen (vgl. Mariotti 2015).

Auch das Erstellen von Untertiteln als Übung für Sprachenlernende wird seit den 1980er-Jahren untersucht und gewinnt mit der stetigen Entwicklung der Informationstechnologie immer mehr an Bedeutung. Das bloße Empfangen audiovisueller Information wurde von vielen Forschenden seit jeher als weniger bedeutungsvoll angesehen, da es eine passive Handlung sei und lange Zeit brauchen würde, um einen positiven Effekt auf die Lernenden auszuüben (vgl. Vanderplank 1988). Das eigenhändige Erstellen von Untertiteln hingegen wurde als aktive Handlung oftmals als hilfreicher für den Lernprozess eingestuft und wurde bereits in Hinblick auf verschiedene Aspekte untersucht, darunter Retention, Erwerb von Lexis, Betonung und Intonation sowie Übersetzungsfähigkeiten. So fand etwa Neves (2004) heraus, dass die verschiedenen Stufen des Untertitelungsprozesses sowohl für die Rezeptions- als auch die Produktionsfähigkeiten von fortgeschrittenen Sprachenlernenden war.

3.3.6 Narration in audiovisuellen Produkten

Der Begriff *Narrativ* wird oft als eine von vielen Kommunikationsformen, darunter auch deskriptive oder argumentative Kommunikation, definiert. Ein narrativer Text zeichnet sich dadurch aus, dass er eine Welt erschafft und diese mit Charakteren füllt. Diese Welt und ihre Charaktere erleben Veränderungen jeglicher Art und Kausalität, und Rezipient:innen können diese Welt mitsamt der Ziele, Pläne, Wirkungszusammenhänge und Motivationen innerhalb dieser interpretieren (vgl. Ryan 2004: 8f.). Ein Narrativ ist also ein Text über eine Welt mit einem Inhalt, also mit Handlungen in temporalen und kausalen Mustern (vgl. Kukkonen 2014: 706). Die Narratologie, also die Wissenschaft der Narration, wurde von Beginn an als transdisziplinär und transmedial gesehen – sowohl Bilder als auch Wörter können eine Narrationsperspektive und somit Narrativität produzieren (vgl. Vandaele 2018: 229).

Vandaele (2018: 229) betont: „[Any] sequentially organized semiotic channel of communication can withhold and gradually disclose information for narrative effect.“ Dieses Vorenthalten von Informationen kann bei Rezipient:innen zwei verschiedene Effekte auslösen – Neugierde oder Überraschung. Das Publikum kann sogar versuchen, die Gründe für das

Vorenthalten zu interpretieren, indem es zum Beispiel auf die Gedankenwelt eines Charakters schließt. Die filmische Narration muss ihren Informationsfluss jedoch nicht rein auf Charaktere fokussieren: Sie kann narrative Perspektive erschaffen, indem Informationen vorenthalten oder gegeben werden, die nicht direkt mit der inneren Welt der Charaktere zu tun haben (vgl. Vandaele 2018: 229f.).

Bei der Translation von audiovisuellen Produkten ist, so Vandaele (vgl. 2018: 230), jedenfalls darauf zu achten, dass die Narrativität und ihre grundlegenden Mechanismen erhalten bleiben. Filme sind so ausgelegt, dass die Rezipient:innen andauernd nach relevanten audiovisuellen Informationen suchen und sich nicht mit irrelevanten visuellen Details aufhalten. Das bedeutet, dass Handlung (verbal, paraverbal und nonverbal) immer auch als Narration fungiert: Die Handlung, die wir sehen, ist die Handlung, die uns mitgeteilt wird (vgl. Vandaele 2018: 232). Filmdialog ist, so Remael (vgl. 2003: 233), nicht nur Dialog, sondern auch Teil des Narrativs. Jegliche Handlung, und damit auch der Dialog, ist sowohl Orientierung der Charaktere in ihrer eigenen Welt als auch Narration für die Zuschauer:innen (vgl. Vandaele 2018: 232).

Diese grundlegenden Überlegungen führen in weiterer Folge zur Frage, wie Rezipient:innen und Translator:innen narrative Relevanz in Filmen finden können. Vandaele (2018) führt dazu in Anlehnung an Pérez-González (2015) und Sternberg (2009) zwei Kriterien an: *narrative clarity* und *temporary opacity*. Narrative clarity bekommen wir durch Hinweise auf Chronologie, Kausalität und Zufall, während temporary opacity nach Interpretationslücken, Hypothesen, verzögerten Antworten und versteckten Informationen sucht. Dies stellt audiovisuelle Translator:innen vor verschiedene Herausforderungen: Aufgrund der jeweils geltenden Konventionen und der Funktion der Untertitel müssen interlinguale Untertitler:innen für relevante Hinweise hauptsächlich auf verbale Informationen zurückgreifen, während SDH-Untertitler:innen zusätzlich paraverbale Aspekte miteinbeziehen müssen (vgl. Vandaele 2018: 232).

Während des Untertitelungsprozesses ist eine Frage vorherrschend: Sind Handlung und Dialog relevant für die Narration sind (etwa für einen Plot Twist zu einem späteren Zeitpunkt)? Sprechakte von Figuren in audiovisuellen Produkten können in inhaltlich relevanten *Dialog* und inhaltlich irrelevante *Konversation* unterteilt werden. In den Untertiteln muss Dialog also (intra- oder interlingual) übersetzt werden, damit narrative Klarheit herrschen kann. Doch auch inhaltlich irrelevante Konversationen können für die Narration relevant werden, wenn sie etwa für humoristische Abwechslung sorgen (vgl. Vandaele 2018: 233). Dialog hat darüber hinaus eine charakterisierende Wirkung, weil Charakterisierung unsere Reaktion auf die Handlung

beeinflusst. Der sogenannte *primacy effect*, die kognitive Salienz des ersten Eindrucks, erlaubt es der Narration, später diese ersten Eindrücke zu widerlegen. Im Gegensatz dazu gibt es den *recency effect*, welcher besagt, dass die letzten Informationen die relevantesten sind. Beide können in audiovisuellen Produkten als Mittel für die Narration eingesetzt werden (vgl. Díaz Cintas & Remael 2007; vgl. Vandaele 2018: 233).

4. Methode und Untersuchungsgegenstand

4.1 Untersuchungsgegenstand: „Behind Her Eyes“ (2021)

Die Serie „Behind Her Eyes“ (deutscher Titel: „Sie weiß von Dir“) basiert auf dem gleichnamigen Roman von Sarah Pinborough. Sie besteht aus sechs Folgen mit jeweils 47 bis 54 Minuten und wurde am 17. Februar 2021 bei Netflix veröffentlicht. Die Dreharbeiten fanden von 2019 bis 2020 unter der Regie von Erik Richter Strand statt; das Drehbuch stammt von Steve Lightfoot und Angela LaManna (vgl. IMDb o.J.).

Die Hauptcharaktere der Serie sind Louise Barnsley (Simona Brown), David Ferguson (Tom Bateman) und dessen Frau Adele Ferguson (Eve Hewson) sowie Rob (Robert Aramayo), dessen Nachname in der Serie nicht genannt wird. Louise ist alleinerziehende Mutter eines Jungen, Adam, und führt ein eher ereignisloses Leben. Sie arbeitet als Assistentin in einer Gemeinschaftspraxis für Psychiater:innen. Als sie David eines Abends in einer Bar kennenlernt und eine Affäre mit ihm eingeht, weiß sie noch nicht, dass er ihr neuer Chef ist. Nach und nach wird die Situation immer verstrickter, und Louise befindet sich inmitten eines Liebesdreiecks mit David und Adele, denn die beiden Frauen freunden sich ebenfalls miteinander an. Sowohl Louise als auch David haben mit Schuldgefühlen zu kämpfen; anfangs können sie ihre Affäre vor Adele verheimlichen, doch gegen Ende der Serie deckt Davids Ehefrau das Geheimnis auf. Im weiteren Serienverlauf werden immer wieder Rückblenden aus Adeles und Davids Leben gezeigt. Als Adele wegen eines psychischen Zusammenbruchs in einer Psychiatrie behandelt wird, lernt sie Rob kennen, mit dem sie eine sehr gute Freundschaft pflegt. Robs Figur ist zu Beginn etwas rätselhaft für die Zuschauer:innen – es scheint ein Geheimnis zu geben, über das David und Adele nicht sprechen (vgl. Framke 2021, Moviepilot o.J.).

Vergangene traumatische Ereignisse aus Louises Leben lösen bei ihr starke Schlafprobleme aus, von denen sie Adele erzählt. Diese schlägt ihr daraufhin eine unkonventionelle Methode vor, welche sie auch aufgrund ihrer Schlafprobleme verwendet, nämlich die Astralprojektion. In der Serie wird Astralprojektion wie folgt definiert:

„Astralprojektion“ ist ein Begriff aus der Esoterik. Er beschreibt eine willentliche außerkörperliche Erfahrung, die die Existenz einer Seele oder eines Bewusstseins voraussetzt, das man 'Astralkörper' nennt. Er unterscheidet sich vom physischen Körper und kann unabhängig von ihm das Universum erkunden. (Folge 5, 31:42-31:58).

Als Louise sich schlussendlich dafür entscheidet, eine Beziehung mit David einzugehen, setzt Adele aus Verzweiflung ihr Haus in Brand und injiziert sich eine große Menge Heroin. Louise möchte versuchen sie zu retten und nutzt die Astralprojektion, um in das Haus zu gelangen,

ohne ihrem Körper dabei Schaden anzurichten. In diesem Moment übernimmt Adele, welche die gleiche Methode benutzt, Louises Körper, wodurch Louises Seele nun in Adeles Körper gefangen ist. Nach dieser Szene finden die Zuseher:innen heraus, wer Adele wirklich ist: Nachdem Rob von Adele die Technik der Astralprojektion gelernt hat, beschließt er, ihren Körper einzunehmen, wie er es gerade mit Louise getan hat. David war all die Jahre also mit Robs Seele verheiratet, und nun wird er mit der vermeintlichen Louise eine Beziehung eingehen. Die Serie endet mit einer emotionalen Szene, denn Louises Sohn Adam bemerkt sofort, dass etwas mit seiner Mutter nicht stimmt, kann die Situation aber nicht mehr beeinflussen (vgl. Framke 2021, Moviepilot o.J.).

4.2 Methodik

Um die ausgewählten Szenen zu analysieren, werden sie zunächst multimodal transkribiert. Die multimodale Transkription wurde ursprünglich für die Filmanalyse entwickelt und wird seit einiger Zeit auch für die Analyse von Translation im Filmkontext verwendet. Sie basiert auf der Segmentierung des Textes in kleinste Bausteine, um eine objektive und empirische Analyse zu ermöglichen. Die auditiven und visuellen Teile des multimodalen Textes werden hierbei in eine Tabelle eingetragen und – abhängig vom Untersuchungszweck – eingeteilt (z.B. können ganze Szenen oder nur einzelne Frames beschrieben werden) (vgl. Pérez-González 2015: 295; Remael & Reviers 2018: 263). Für meine Analyse möchte ich die folgenden Faktoren in meine Transkription einbeziehen: Nummerierung (der Beispiele), sichtbare Handlung, Transkription der Musik und der Geräusche, Transkription des Dialogs und Transkription der SDH-Untertitel. Bei Szenen, die keinen (relevanten) Dialog enthalten, entfällt folglich auch die Beschreibung. Die Unterpunkte Dialog und SDH-Untertitelung wiederum enthalten alle drei Sprachfassungen, damit diese so gut wie möglich miteinander verglichen werden können. Da viele Beispiele aus mehr als nur einem Untertitel bestehen, werden die einzelnen Untertitel in der Transkription durch Unterteilung sichtbar gemacht.

Anschließend werden die Szenen mithilfe des Analysemodells für audiovisuelle Texte nach Chaume (2004), welches in den folgenden Absätzen näher beschrieben wird, untersucht. Chaume schlägt in seinem Artikel „Film Studies and Translation Studies: Two Disciplines at Stake in Audiovisual Translation“ (2004) einen Ansatz vor, welcher film- und kommunikationswissenschaftliche Aspekte mit translationswissenschaftlichen Aspekten vereint. Ein Analyserahmen, der den ganzen Bedeutungsinhalt eines Textes untersuchen

können soll, solle, so Chaume (vgl. 2004: 16), sowohl externe Faktoren des Textes als auch generelle Übersetzungsprobleme beinhalten.

Basierend auf einer multimodalen Kommunikationstheorie beschreibt Chaume (2004), dass audiovisuelle Texte semiotische Konstrukte seien, dessen verschiedene bedeutungstragende Codes gleichzeitig eine Bedeutung herstellen. Ein Film sei eine Reihe kodifizierter Zeichen in einer bestimmten syntaktischen Reihenfolge. Seine Typologie, Organisation und die Bedeutung seiner Elemente werden zu einer semantischen Struktur, welche die Zuseher:innen entziffern, um die Bedeutung des Textes zu verstehen. Für die Translator:innen audiovisueller Texte sei also am wichtigsten, die Funktion aller Codes und das Auftreten aller Zeichen (sprachlich und nicht-sprachlich) in einer Übersetzung zu kennen (vgl. Chaume 2004: 16f.).

Angelehnt an Ansätze aus der Übersetzungswissenschaft, Diskursanalyse und Film- und Kommunikationswissenschaften basiert dieses Modell auf verschiedenen Codes, deren Zusammenspiel einen audiovisuellen Text ausmacht. Im Anschluss folgt eine kurze Beschreibung aller Codes, deren Entsprechung nach Stöckls Theorie der Multimodalität (2004) und deren Bedeutung für die Untertitelung für Hörgeschädigte:

Der *linguistic code* ist, so Chaume (2004), nicht unbedingt eines der Alleinstellungsmerkmal audiovisueller Texte, da auch andere Formen der Translation diesen Code beinhalten. Dieser Code deckt sich mit dem Modus Sprache nach Stöckl (2004). Was aber den linguistic code in audiovisuellen Texten so einzigartig macht, ist, dass es Texte sind, die geschrieben wurden, aber nicht sofort so wirken sollen – sie sollen in den meisten Fällen eher wie spontane, authentische Äußerungen klingen. Dafür stehen sowohl den Drehbuchautor:innen des Ausgangsproduktes als auch den Translator:innen mehrere Strategien zur Verfügung, etwa Interjektionen, Auslassungen und Assimilationen („coulda“ bzw. „gimme“) (vgl. Chaume 2004: 17). Für die SDH-Untertitelung kann gerade Alltagssprache sehr herausfordernd sein, da (je nach Richtlinien) die Sprache zwar alltäglich sein soll, aber nicht so kompliziert, dass Hörgeschädigte sie aufgrund etwaiger niedrigerer sprachlicher Kompetenz (s. Kapitel 3.2.3 zur Heterogenität der Zielgruppe) eventuell nicht mehr verstehen könnten.

Der *paralinguistic code* beschreibt paraverbale Merkmale, zum Beispiel Lautstärke, Sprechpausen oder auch Dialekte von Sprecher:innen. Auch diese Besonderheiten finden sich bei Stöckl (2004) im Modus Sprache wieder. Wie bereits in Kapitel 1.1.1.2 beschrieben, können gerade paraverbale Merkmale audiovisuelle Translator:innen vor große Herausforderungen stellen, vor allem, wenn keine hundertprozentige Entsprechung in der Zielsprachlichen Fassung vorhanden ist (z.B. bei Dialekten oder Akzenten). Paraverbale Besonderheiten können in SDH-

Untertiteln angezeigt werden, etwa durch Klammern ([leise:]) oder auch durch Symbole bzw. Abkürzungen, etwa „--“ bei Unterbrechungen (vgl. Chaume 2004: 17).

Der *musical and special effects code* – also die Modi Musik und Ton nach Stöckl (2004) – vereint den Soundtrack des audiovisuellen Produktes und die Spezialeffekte. Es kann sich hierbei also sowohl um Lieder handeln, die im audiovisuellen Text vorkommen, als auch um Effekte wie Applaus, das Schließen einer Türe oder einen Knall (vgl. Chaume 2004: 18). Diese diegetischen Elemente eines audiovisuellen Produkts sind ein Hauptbestandteil von SDH-Untertiteln und sollten im besten Fall in den Untertiteln reproduziert werden, damit hörgeschädigte Zuseher:innen das audiovisuelle Produkt in gleichem Maße verstehen können wie das hörende Publikum.

Als *sound arrangement code* beschreibt Chaume (2004: 18) spezifische Merkmale der Tonspur, etwa, ob der zu hörende Ton diegetisch oder nicht-diegetisch ist und ob er on oder off screen produziert wird. Dieser Code entspricht dem Modus Ton nach Stöckl (2004) und ist ebenfalls für die Untertitelung für Hörgeschädigte relevant, da zum Beispiel Farben oder auch Kursivsetzung für Off-Sprecher:innen eingesetzt werden können. In einigen Fällen wird dieses Merkmal auch durch den Ausdruck (off:) oder (on:) angezeigt.

Der *iconographic code* ist, so Chaume (2004: 18f.), der relevanteste visuell vermittelte Code. Er korreliert wie alle nachfolgenden Codes mit dem Modus Bild nach Stöckl (2004) und beschreibt die bildlich dargestellten Elemente eines audiovisuellen Produktes. In den meisten Fällen ist es nicht nötig, diese Elemente in SDH-Untertiteln explizit darzustellen, da die Rezipient:innen diese selbst wahrnehmen können.

Der *photographic code* befasst sich mit Veränderungen der Lichtverhältnisse, der Perspektive oder der Verwendung von Farbe (vgl. Chaume 2004: 19). Dieser Code kann ebenfalls einen Einfluss auf die SDH-Untertitelung haben – zum Beispiel, wenn ein:e Sprecher:in aufgrund schlechter Lichtverhältnisse nicht zu sehen ist, seine/ihre Stimme aber zu hören ist. Auch die Erwähnung von Farben in der Tonspur und die gleichzeitige Erscheinung dieser Farbe am Bildschirm kann eine diegetische Funktion erfüllen, weshalb auf solche Merkmale zu achten ist.

Der *planning code* bezieht sich auf Kameraeinstellungen und Veränderungen der Perspektive, der Farbeinstellungen oder des Lichts und ist daher weniger relevant für die Untertitelung für Hörgeschädigte. In der interlingualen Untertitelung jedoch kann dieser Code sehr wichtig sein, wenn zum Beispiel ausgangssprachliche Inhalte durch eine Nahaufnahme auf dem Bildschirm zu lesen sind und für das Zielpublikum übersetzt werden müssen (vgl. Chaume 2004: 19f.).

Der *mobility code* beschreibt die Bewegung von Charakteren in einem audiovisuellen Produkt, also proxemische und kinetische Zeichen oder auch die Mundbewegungen der Charaktere (vgl. Chaume 2004: 20). Es kann für die SDH-Untertitelung relevant sein, welche Charaktere gerade im Bild zu sehen sind und wie weit diese von der Kamera entfernt sind. Audiovisuelle Translator:innen müssen diese Aspekte eventuell in die Untertitelung integrieren. Auch Synchronität spielt bei diesem Code eine Rolle – sowohl in Bezug auf Synchronität zwischen Ton und Bild (Lippensynchronität) als auch in Bezug auf die physische Nähe zweier Charaktere, welche miteinander sprechen.

Der *graphic code* stellt geschriebene Sprache in audiovisuellen Produkten dar, worunter auch die Untertitelung fällt. Andere Ausprägungsarten sind Titel und Zwischentitel sowie Texte (zum Beispiel SMS-Nachrichten). Dieser Code kann aufgrund von Richtlinien und auch aufgrund anderer Codes variieren, zum Beispiel, wenn ein Untertitel an den oberen Bildrand geschoben wird, um keine Informationen auf dem Bildschirm zu verdecken, oder auch, wenn Synchronität zwischen Untertitel und Ton hergestellt werden soll (vgl. Chaume 2004: 20).

Der *syntactic code* bezieht sich auf die Beziehung der einzelnen Szenen zueinander und ihre Bedeutung für den audiovisuellen Text, also auf den Schnitt. Die Schnittfolge kann den Translator:innen dabei helfen, Zusammenhänge besser zu verstehen und die Untertitelung dementsprechend anzupassen. Dabei ist es aber wichtig, dass die Untertitel das gleiche Maß an Informationen preisgeben wie der Ausgangstext. Außerdem wird bei allen Untertitelungsarten darauf geachtet, die Untertitel so synchron wie möglich mit der Tonspur erscheinen zu lassen und dabei auch Schnitte zu beachten, damit ein Untertitel nicht zu knapp vor oder nach einem Schnitt endet oder zu lange in einer neuen Szene zu sehen ist (vgl. Chaume 2004: 20f.).

All diese Codes können also verwendet werden, um ein audiovisuelles, multimodales Produkt zu beschreiben – und genauso können sie auch besondere Herausforderungen für die Untertitler:innen darstellen. In meiner Analyse möchte ich jenen Codes Vorrang geben, die für die Untertitelung im jeweiligen Fall relevant sind – in den meisten Fällen sind das die sprach-, musik- und geräuschbezogenen Codes, wobei auch andere selbstverständlich eine Rolle spielen können.

4.3 Ziel der Analyse

Ziel der Untersuchung ist es, herauszufinden, ob und inwiefern die Gemeinsamkeiten und Unterschiede in den Situations- und Kontextangaben der deutschen, englischen und französischen Sprachfassungen Einfluss auf die filmische Narration haben. Dazu werde ich die

Serie in ihren einzelnen sprachlichen Fassungen ansehen und möglichst repräsentative Beispiele für die Analyse heranziehen. Anschließend möchte ich die gewählten Beispiele transkribieren und anhand des Modells nach Chaume (2004) untersuchen, inwiefern Situations- und Kontextangaben die filmische Narration in diesen bestimmten Szenen beeinflussen.

Es geht bei der Untersuchung sowohl um die Unterschiede zwischen den einzelnen Sprachfassungen – etwa, ob eine Version mehr Informationen preisgibt als die andere – als auch um den Einfluss der Situations- und Kontextangaben einer bestimmten Sprachfassung. Ich erhoffe mir aussagekräftige Ergebnisse, die im besten Fall auch auf andere audiovisuelle Produkte übertragen werden könnten oder zumindest damit vergleichbar sind.

5. Untersuchung der Situations- und Kontextangaben in „Behind Her Eyes“

5.1 Musik

5.1.1 Musikbeschreibungen

5.1.1.1 Titelmelodie

Einen großen Unterschied in Hinblick auf die Musikbeschreibungen findet man schon bei der Titelmelodie, die in allen Episoden während des Intros abgespielt wird (z.B. Folge 1, 7:18-7:26).

Nummer	Sichtbare Handlung	Transkription Musik, Geräusche
1	Der Titel „Behind Her Eyes“ erscheint nach und nach vor einer Animation eines düsteren Waldes; eine Lichtquelle leuchtet kurz hinter den Bäumen auf und erlischt wieder.	Schaurige, düstere Musik

Nummer	Transkription SDH		
1	Englisch	Deutsch	Französisch
	---	[geheimnisvolle Titelmelodie]	---

Wie an diesem Beispiel sichtbar wird, erfolgt nur in der deutschen Version eine Beschreibung der Titelmelodie – diese Strategie zur Untertitelung bzw. Nicht-Untertitelung dieses Aspekts des *musical and special effect codes* nach Chaume (2004) zieht sich durch alle Folgen der Serie.

Die Titelmelodie kann einer der Hauptindikatoren für die Richtung sein, in die die Serie sich bewegen wird – aus ihr können viele Informationen abgelesen werden. Ich würde die Melodie als düster und schaurig beschreiben, was sowohl zur Animation im Hintergrund passt, die ebenfalls sehr düster ist, als auch zur allgemeinen Stimmung der Serie (Thriller/Drama). Der Wald im Intro steht für den Wald nahe Adeles Anwesen, in dem Rob getötet worden ist (was die Zuseher:innen erst zu einem späteren Zeitpunkt herausfinden), und ist daher einer der Hauptschauplätze der Serie. Die Buchstaben „Behind Her Eyes“ sind zu Beginn nicht sichtbar; erst nach und nach werden sie von orangener Farbe ausgefüllt, die wie Blut von ihnen herablaufen. Gleichzeitig wird auch der Wald erhellt, und die düstere Musik wird lauter. Das ganze Erscheinungsbild des Intros kommt also erst nach und nach zum Vorschein. Umso

wichtiger wäre es für das Verständnis dieses kurzen Clips, die Musik auch in den Untertiteln zu erwähnen.

In dieser Form können nur die deutschsprachigen Zuseher:innen die Titelmelodie nachvollziehen, da sie in den anderen Sprachfassungen nicht erwähnt wird. Der Modus Musik hat großen Einfluss auf die Rezeption von Charakteren und Settings (vgl. Pérez-González 2015) und ist daher ein sehr wichtiger Bestandteil für das Verständnis von audiovisuellen Produkten. Fehlt ein solcher integraler Baustein des multimodalen Textes, kann das dazu führen, dass die Rezipient:innen die Zusammenhänge im Intro der Serie nicht im gleichen Maß nachvollziehen können wie die hörenden Rezipient:innen.

5.1.1.2 Generische Musikbeschreibungen

Auch die generischen Musikbeschreibungen (also jene, die sich nicht auf einen bestimmten Song beziehen) variieren von Sprache zu Sprache. In diesem Abschnitt möchte ich ebenfalls ein Beispiel aus Folge 2 (9:35-11:15) analysieren, in dem die Hintergrundmusikbeschreibungen sehr gut veranschaulicht werden. Ähnlich wie im zuvor analysierten Beispiel geht es in dieser Szene auch um den Wald als Schauplatz der Serie. (Aus Platzgründen wird nicht die gesamte Szene transkribiert, sondern nur die Musikbeschreibungen.)

Nummer	Sichtbare Handlung	Transkription Musik, Geräusche
2	Immer wieder werden verschiedene Szenen gezeigt: • Adele malt eine düstere Waldlandschaft auf ihre Schlafzimmerwand. • Ein dunkler Wald, immer wieder sieht man einen Brunnen. Am Ende schwenkt die Kamera auf den Grund des Brunnens, wo man Davids Uhr liegen sieht.	Geheimnisvolle Musik, die sich gegen Ende der Szene steigert und dann abrupt endet

Nummer	Transkription SDH		
2	Englisch	Deutsch	Französisch
	[tense music playing]	[spannungsvolle Musik]	[musique sombre]
		[bedrohliche Musik]	
		[Musik wird intensiver]	

		[Musik wird leiser und düsterer]	[musique lugubre]
		[Musik wird intensiver, unheilvoller]	
		[Musik verklingt abrupt]	

Der Wald ist ein wiederkehrendes Thema in dieser Serie: Noch bevor die Zuseher:innen herausfinden, welche tiefere Bedeutung der Wald für die Hauptcharaktere David und Adele hat, erfolgt in dieser Szene eine weitere Referenz auf ihn. Darüber hinaus werden diese Szenen mit düsterer Musik untermalt – dies entspricht dem *musical and special effects code* nach Chaume (2004). Wie anhand dieses Beispiels eindeutig zu erkennen ist, werden die Veränderungen der Musik im Deutschen beschrieben, zum Beispiel, wenn Emotion oder Lautstärke verändert werden oder die Musik endet. Die filmische Narration, die sich durch die Musik vollzieht, ist also für die deutschsprachigen Rezipient:innen besser nachvollziehbar als für die englisch- und französischsprachigen. Hinzu kommt, dass in den beiden Sprachen während der ganzen Szene, die ja beinahe zwei Minuten andauert, nur ein einziges Mal auf die Musik verwiesen wird und dazwischen auch andere Geräuschbeschreibungen erfolgen. Es ist also meiner Meinung nach in diesen Sprachen nicht unbedingt klar, dass im Hintergrund noch immer Musik zu hören ist, wenn nicht erneut auf sie verwiesen wird.

Die hier beschriebenen Untertitelungsstrategien ziehen sich durch die gesamte Serie – die Musik wird in den französischsprachigen Untertiteln oftmals gar nicht erwähnt, in den englisch- und deutschsprachigen in den meisten Fällen schon. Darüber hinaus erfolgen im Deutschen oft nähere Beschreibungen zur Musik, etwa, wenn diese lauter wird oder die Stimmung sich verändert. Es kommen unter anderem Untertitel wie [Musik wird leiser und düsterer] oder [Musik verklingt] zum Einsatz, wie auch im obigen Beispiel zu sehen ist. Das hat zur Folge, dass die Veränderungen in der Musik, die auch immer diegetische Funktion haben, für die deutschsprachigen Rezipient:innen wahrscheinlich am besten nachvollziehbar ist, während dem englisch- und französischsprachigen Publikum keine derartigen Informationen zur Verfügung stehen.

Zusammenfassend für die Unterpunkte der Titelmelodie und der generischen Musikbeschreibungen lässt sich sagen, dass den deutschsprachigen hörgeschädigten Rezipient:innen die genauesten Musikbeschreibungen zur Verfügung stehen. In den anderen beiden Sprachen wird weder die Titelmelodie beschrieben, noch erfolgen genauere Beschreibungen der Hintergrundmusik oder deren Veränderungen. Die Emotion und folglich

auch die Narration, die durch die Musik transportiert werden soll, ist für die deutschsprachigen Zuseher:innen also nachvollziehbarer als für die englisch- oder französischsprachigen. Im Bereich der interlingualen audiovisuellen Translation haben vielzählige Studien gezeigt, dass hörende Rezipient:innen die Stimmung eines Songs auch dann einschätzen können, wenn sie den Liedtext nicht verstehen (vgl. Pérez-Gonzales 2015: 208f.). Diese Erkenntnis ließe sich wahrscheinlich auch auf die hier vorliegende Situation anwenden: Wenn hörgeschädigte Personen den Text eines Liedes schon nicht hören können bzw. gar kein Liedtext existiert, dann könnten sie durch Musikbeschreibungen wenigstens die Stimmung erahnen. Das würde sicherstellen, dass sie den *musical and special effects code*, der durch die Hintergrundmusik dargestellt wird, (im besten Fall) in gleichem Maße nachvollziehen können wie hörende Personen.

5.1.2 Musiktitel und -texte

Hinsichtlich der Beschreibung von Musiktiteln und Songtexten bestehen in den einzelnen Sprachräumen verschiedene Herangehensweisen, welche auch in den Untertiteln und in den sprachbezogenen Richtlinien abzulesen sind. Während die englischen und französischen SDH-Untertitel nur den Songtitel erwähnen – zum Beispiel am Ende von Episode 3: [„Rocking Horse“ playing] oder [chanson „Rocking Horse“]) –, werden in den deutschen Untertiteln auch die Interpreten genannt: [Musik: „Rocking Horse“ von Kelli Ali] (Episode 3, 48:17-49:15). Deutschsprachige Zuseher:innen können also besser nachvollziehen, welcher Song gespielt wird, weil ihnen mehr Informationen zur Verfügung stehen. Der *music and special effects code* nach Chaume (2004) wird also auf unterschiedliche Art und Weise dargestellt.

Die Netflix-Richtlinien aller drei Sprachen geben vor, dass Musiktexte, die in der gleichen Sprache verfügbar sind wie die Audiospur, Untertitelt werden sollen. Da „Behind Her Eyes“ eine britische Serie ist und die englischen Songtexte beim Untertitelungsprozess nicht an die Zielsprache angepasst werden, werden sie in den meisten Fällen auch nur in den englischen SDH-Untertiteln transkribiert. Das oben genannte Beispiel („Rocking Horse“ von Kelli Ali) stellt jedoch einen Spezialfall dar, da der Liedtext sowohl in den englischen Untertiteln transkribiert wird, als auch in den deutschsprachigen Untertiteln eine übersetzte Version zu finden ist. In der französischen Sprachfassung wird der Liedtext weder auf Englisch transkribiert, noch übersetzt. Transkribiert sieht diese Szene wie folgt aus:

Nummer	Sichtbare Handlung	Transkription Musik, Geräusche
3	Schwarzer Hintergrund, Nachspann der Episode in weißen Buchstaben	Ruhige Musik, singende Frauenstimme

Nummer	Transkription Dialog		
3	Englisch	Deutsch	Französisch
	You must set your demons free. I am you and you are me. Chasing dreams around the bend. Going out my mind again. [...]	You must set your demons free. I am you and you are me. Chasing dreams around the bend. Going out my mind again. [...]	You must set your demons free. I am you and you are me. Chasing dreams around the bend. Going out my mind again. [...]

Nummer	Transkription SDH		
3	Englisch	Deutsch	Französisch
	["Rocking Horse" playing]	[Musik: "Rocking Horse" von Kelli Ali]	[chanson "Rocking Horse"]
	♪ You must set your demons free ♪	♪ Setz deine Dämonen frei für mich ♪	
	♪ I am you and you are me ♪	♪ Ich bin du, und du bist ich ♪	
	♪ Chasing dreams around the bend ♪	♪ Ich jage meinen Träumen nach ♪	
	♪ Going out my mind again ♪ [...]	♪ Mein Verstand, der liegt ganz brach ♪ [...]	

Obwohl dieser Song am Ende einer Folge vorkommt, ist er doch sehr relevant für den Plot und sollte daher auf jeden Fall transkribiert werden. Die vorhergehende Szene zeigt Adele und David bei sich zu Hause, und Adele spielt darauf an, dass David sie nicht verlassen könne, weil sie in der Vergangenheit viel miteinander erlebt hätten und viele Geheimnisse teilen würden. Zu diesem Zeitpunkt wissen die Zuseher:innen noch nicht, wovon Adele spricht. Der Songtext bezieht sich eindeutig auf diese Situation; darauf, dass David und Adele ihre Dämonen aus der Vergangenheit freisetzen sollen, von denen sie gerade gesprochen haben. Obwohl in keiner der Sprachen eine genauere Musikbeschreibung erfolgt, die die Stimmung näher beschreibt, ist diese durch den Liedtext abzulesen. Er beschreibt Verzweiflung und Verwirrung – beides Gefühle, auf die Adele und David in ihrer Diskussion eingehen.

Die englisch- und deutschsprachigen Zuseher:innen, die den Originaltext bzw. die Übersetzung ablesen können, können also zumindest nachvollziehen, welche Stimmung der Liedtext ausstrahlt⁶. In der französischen Sprachfassung gibt es jedoch keine Möglichkeit, das unmittelbar aus den Untertiteln herauszufinden – für sie könnte „Rocking Chair“ jeder beliebige Song sein, vor allem, weil er in den End-Credits abgespielt wird und so auch keine Hinweise aus dem Kontext abzulesen sind. Das könnte zu einer ähnlichen Situation wie bei Desillas Untersuchung (2009) führen, in der das Publikum der Zielsprache nicht erklären kann, wieso ein bestimmtes Lied in einer bestimmten Situation gewählt wurde. In weiterer Folge könnte das Nichtübersetzen bzw. Nicht-Transkribieren von Liedtexten sogar dazu führen, dass der multimodale Text nicht oder nur oberflächlich verstanden wird (vgl. Di Giovanni 2008; Pérez-Gonzales 2015).

Ich möchte hier erneut betonen, dass das oben genannte Beispiel in der deutschen Sprachfassung einen Spezialfall darstellt – in den meisten anderen Fällen wurden die Songtexte nicht übersetzt. Zu Beginn von Folge eins etwa ist lebhafte Hintergrundmusik („Never Forget You“ von Noisettes) zu hören, die Louises Gemütszustand unterstreichen soll: Als sie sich für ein Treffen mit einer Freundin fertig macht, tanzt sie zur Musik in ihrem Schlafzimmer und betrachtet sich selbst im Spiegel. Der Song im Hintergrund, in dem eine Frau über das Ausgehen und das Trinken von Alkohol singt, und außerdem jemandem verspricht, diese Person niemals zu vergessen („I’ll never forget you“), könnte eine Anspielung darauf sein, was Louise an diesem Abend noch erleben würde. Der Text wird jedoch nur in den englischsprachigen Untertiteln transkribiert; in der deutschen Version werden Songtitel und Interpret erwähnt, in der französischsprachigen Version nur der Titel (ähnlich wie im oben genannten Beispiel), aber der Text wird nicht transkribiert. Während englischsprachige Rezipient:innen die Anspielung also verstehen können, haben deutsch- und französischsprachige Zuseher:innen diese Möglichkeit nicht, was die gleichen Folgen wie im obigen Beispiel haben könnte, nämlich kein bzw. nur oberflächliches Verständnis des multimodalen Textes (vgl. Desilla 2009; Di Giovanni 2008).

⁶ Die deutsche Übersetzung achtet sogar auf das Ausgangsformat, da sie das Reimschema des Originaltextes beibehält und den Text eher an dieses anpasst, als ihn wörtlich zu übersetzen. Dadurch entsteht allerdings ein geringer Bedeutungsunterschied im Vergleich zum Original.

5.2 Geräusche

Ein sehr aussagekräftiges Beispiel, anhand dessen die Geräuschbeschreibungen in der Serie dargestellt werden können, findet sich in Folge 3 (1:44-3:09), in der für relativ lange Zeit wenig gesprochen wird, aber viel Information über Musik und Hintergrundgeräusche transportiert wird.

Nummer	Sichtbare Handlung	Transkription Musik, Geräusche
4	<u>Szene 1:</u> David verlässt früh morgens das Haus und versucht dabei leise zu sein, um Adele nicht aufzuwecken. <u>Szene 2:</u> Die Szene wechselt ins Büro, wo Louise die Kaffeemaschine bedient und sich anschließend in Davids Büro setzt. Als sie aufsteht, betritt David plötzlich sein Büro.	<u>Szene 1:</u> Anfangs keine Musik, Schritte, Stufe knarzt, Schlüssel klimpern, Adele atmet laut, Türe öffnet sich und fällt zu. Melancholische Klaviermusik beginnt, Vogelzwitschern, Hundebellen. <u>Szene 2:</u> Melancholische Musik spielt weiter, Klappern der Kaffeemaschine, Louise seufzt.

Nummer	Transkription Dialog		
4	Englisch	Deutsch	Französisch
	[Nachbarin] Good morning.	[Nachbarin] Guten Morgen.	[Nachbarin] Bonjour.

Nummer	Transkription SDH		
4	Englisch	Deutsch	Französisch
	[car engines roar in distance]		
	[steps creaking]	[Stufe knarrt]	[grincement d'escaliers]
	[keys clinking]	[Schlüssel klimpern]	
	[heavy breathing]		
	[door opens]		
	[door closes]	[Tür fällt zu]	
	[melancholy music playing]	[mysteriöse Musik]	[musique mélancolique]
	[birds chirping]	[Vögel zwitschern, Hund bellt]	
	Good morning.	Guten Morgen.	[femme] Bonjour.
			[musique mélancolique]
	[exhales]	[seufzt]	[souponir]
		[Musik wird lauter]	
		[Musik verklingt]	

In diesen Szenen sind gleich mehrere audiovisuelle Codes nach Chaume (2004) vorhanden: Die Musik und die hörbaren Geräusche etwa entsprechen dem *music and special effects code*, die sprachliche Handlung dem *linguistic code*, und die paraverbalen Merkmale der Charaktere (z.B. das Seufzen und Atmen) werden unter dem *paralinguistic code* zusammengefasst. Diese Merkmale werden in Kapitel 5.3.3 genauer analysiert. Auch der *sound arrangement code* spielt hier eine wichtige Rolle, denn er legt fest, welche Teile der Tonspur diegetisch bzw. nicht diegetisch sind und daher untertitelt werden sollen (oder nicht).

In diesen Szenen gibt es, wie zuvor erwähnt, für relativ lange Zeit (verglichen mit dem Rest der Serie) keinen Dialog, dafür werden umso mehr handlungsrelevante Informationen über die Tonspur vermittelt. Da in beiden Szenen die gleiche Hintergrundmusik spielt, ist die Verbindung zwischen ihnen eindeutig erkennbar. Auch die gehörlosen Zuseher:innen können auf diese Information zugreifen, da sie in allen Sprachversionen erwähnt wird, jedoch wird die Musik unterschiedlich stark betont, was sich mit den in Kapitel 5.1 identifizierten Strategien zur Musikbeschreibung deckt.

Auffallend ist vor allem die Anzahl an Untertiteln in den jeweiligen Sprachfassungen – die englische Version (12 UT) weist vor der deutschen (10 UT) und französischen (7) die meisten Untertitel auf. Das wirkt sich in weiterer Folge auch auf die Genauigkeit der Tonbeschreibung aus, da in der englischen und deutschen Version eine größere Anzahl an Informationen vermittelt werden kann. Wie bereits zu Beginn beschrieben, möchte sich David zu Beginn der Szene aus dem Haus schleichen, ohne seine Frau Adele aufzuwecken – er versucht daher, sich so leise wie möglich im Haus zu bewegen. Als er jedoch die Stufen hinuntergeht, knarzt eine von ihnen; abgesehen davon ist zu dieser Zeit kein anderer Ton zu hören. Das Knarzen der Stufe hat hier also durchaus eine diegetische Funktion, die auch in allen drei Sprachen in den Untertiteln reproduziert wird.

Dass David möglichst leise sein möchte, ist schon an seiner Körpersprache abzulesen – er ist sehr aufmerksam und bewegt sich nur langsam. Es ist also eindeutig, dass er unter allen Umständen vermeiden möchte, dass Adele ihn hört; das laute Knarzen der Stufen könnte seinen Plan durchkreuzen und Adele aufwecken. Es ist daher umso wichtiger, dieses Geräusch zu erwähnen, da es eine wichtige Funktion für die Narration erfüllt. Das wichtigste Element in dieser Szene wurde also in allen drei Sprachen beschrieben – in der englischen und deutschen Version werden darüber hinaus auch die Hintergrundgeräusche genauer erwähnt, zum Beispiel das Klimpern der Schlüssel, welches Davids Absicht, das Haus leise zu verlassen, auch im Wege stehen könnte.

In diesem Beispiel sind in allen drei Sprachen die meiner Meinung nach wichtigsten – handlungsrelevanten – Beschreibungen enthalten. Die gehörlosen Zuseher:innen jeder Version können die diegetischen Informationen also anhand der Untertitel ablesen. In diesem Fall sind das in Szene eins das Knarzen der Treppe, Adeles Seufzen und auch die Musik; in Szene zwei sind es die paraverbalen Merkmale, also Louises Ausatmen und das Aufschrecken von Louise und David. Das bedeutet, dass die Zuseher:innen auch die Narration nachvollziehen können müssten – der Aspekt der *narrative clarity* nach Vandaele (2018) sollte jedenfalls in allen drei Sprachfassungen erfüllt sein. Jedoch können die englisch- und deutschsprachigen Rezipient:innen die multimodale Welt der Charaktere wahrscheinlich besser verstehen, da ihnen mehr Informationen zu den nicht-diegetischen Geschehnissen zur Verfügung stehen, etwa das Hundebellen oder das Zwitschern der Vögel. Diese Merkmale sind nicht zwingend relevant für die Handlung, machen die Welt, in der die Charaktere sich bewegen, jedoch lebendiger und sind daher ebenfalls ein wichtiger Bestandteil des multimodalen Produktes.

5.3 Charaktere

5.3.1 Akzente

5.3.1.1 Schottischer Akzent

Um den Aspekt der Akzente in „Behind Her Eyes“ zu untersuchen, möchte ich zunächst eine Szene analysieren, die gleich zu Beginn der ersten Folge auftritt. Die Hauptpersonen Louise und David treffen zum ersten Mal aufeinander und führen anschließend in einer Bar ein Gespräch über den Whiskey, den David bestellt hat (Folge 1, 05:09-05:15).

Nummer	Sichtbare Handlung	Transkription Musik, Geräusche
5	David und Louise unterhalten sich in der Bar, die Kamera schwenkt (in Großaufnahme) zwischen den beiden hin und her.	Hintergrundkulisse der Bar (Gespräche, Lachen), Musik

Nummer	Transkription Dialog		
5	Englisch	Deutsch	Französisch
	[David] It's Scottish. [Louise] Like you. I like your accent.	[David] Aus Schottland. [Louise] Mhm, wie Sie. Nettes Völkchen, die Schotten.	[David] C'est un Écossais. [Louise] Comme vous. J'adore les Écossais.

Nummer	Transkription SDH		
5	Englisch	Deutsch	Französisch
	It's Scottish.	Aus Schottland.	C'est un Écossais.
	Like you.	Mhm. Wie Sie.	Comme vous.
	I like your accent.	Nettes Völkchen, die Schotten.	J'adore les Écossais.

Bei dieser Szene gibt es gleich mehrere translationsrelevante Herausforderungen: Zum einen spricht David in der englischen Originalfassung mit einem schottischen Akzent, welcher (wie üblich) in den synchronisierten Fassungen nicht zu hören ist und daher auch für die hörenden Zuseher:innen nicht wirklich aus dem Kontext nachvollziehbar ist. Gleichzeitig ist die Information, dass David mit Akzent spricht, auch nicht in den SDH-Untertitelungen (egal welcher Sprachfassung) enthalten, weshalb auch das gehörlose Publikum Louises Aussage nicht unmittelbar nachvollziehen kann. Die Äußerung ist also nur für eine bestimmte Gruppe an Rezipient:innen verständlich, nämlich hörende englischsprachige Menschen.

Dieses Beispiel zeigt eine große Herausforderung im *paralinguistic code* nach Chaume (2004): Da Akzente und andere paraverbale Merkmale – wie bereits in Kapitel 1 beschrieben – großen Einfluss auf die Charakterisierung von Figuren in multimodalen Produkten haben können (vgl. Queen 2004, Bosseaux 2008), stellt sich also die Frage, inwiefern sowohl gehörlose englischsprachige als auch hörende nicht-englischsprachige Zuseher:innen auf diese Information zugreifen sollen. Um den paralinguistic code der englischen SDH-Untertitel an die Tonspur anzugleichen, könnte man etwa die Kontextangabe [Scottish accent:] gleich zu Beginn des Gesprächs zwischen David und Louise einfügen.

In den anderen beiden Sprachen sowie jeder anderen Synchronfassung ist dieses Problem eher ein inter- als ein intralinguales, da die Untertitler:innen in den meisten Fällen in der Postproduktion mit intralingualem Tonmaterial arbeiten und daher keinen Einfluss auf die Synchronisation haben. In diesem Fall läge es also eher an den Ersteller:innen der Synchronfassung, eine geeignete Lösung zu finden, und zum Beispiel im Gespräch Davids Herkunft anstelle seines Akzents zu erwähnen, weil dieser ja nicht zu hören ist. In beiden Synchronfassungen wird zwar auf die Schotten im generellen Sinn hingewiesen („Nettes Völkchen, die Schotten“ bzw. „J'adore les Écossais“), aber es wird nicht explizit erwähnt, wieso Louise David zu dieser Gruppe zählt.

Obwohl gehörlose Rezipient:innen Situations- und Kontextinformationen durch eigene Erfahrungen oder Weltwissen oftmals auch ohne Untertitelungen verstehen können (vgl.

Remael & Reviers 2018; Tiittula & Hirvonen 2019), ist es hier meiner Meinung nach schon notwendig, Davids Akzent zu erwähnen, weil er zu einem zentralen Bestandteil des Dialogs – und damit zu einem diegetischen Teil der Narration – geworden ist und über die „bloße“ Charakterisierung seiner Person hinausgeht (was laut Netflix-Richtlinien ebenfalls ein Grund wäre, diesen zu erwähnen). Im besten Fall können sich jedoch nur englischsprachige gehörlose Menschen diese Information durch Weltwissen erklären; die Rezipient:innen anderer Sprachfassungen haben diese Möglichkeit nicht. Außerdem erwähnt Louise im englischsprachigen Original sogar noch einmal, woher genau Davids Akzent kommt, sodass auch Menschen, die ihn zuvor noch nicht zuordnen konnten, danach als schottisch verorten können.

In dieser Form ist der Akzent also nur für eine Rezipient:innengruppe ohne Weiteres zugänglich – für nicht-englischsprachige und gehörlose Zuseher:innen gibt es keine Möglichkeit, über die Untertitel nachzuvollziehen, wieso Louise sagt, dass der Whiskey – wie David – aus Schottland stammt. Eine derartige Situation, in der ein wichtiger Aspekt eines Charakters nicht aus der Untertitelung abzulesen ist, könnte zu einem höheren mentalen Aufwand für die Rezipient:innen führen, was dem eigentlichen Sinn der Barrierefreiheit entgegenstehen würde. Das könnte in weiterer Folge die Auswirkung haben, dass das hörgeschädigte Publikum die filmische Narration (oder zumindest einen Teil davon) dieses audiovisuellen Produktes nicht nachvollziehen kann, da keine *narrative clarity* (vgl. Vandaele 2018) gegeben ist.

5.3.1.2 Französischer Akzent

Ein Beispiel für eine (meiner Meinung nach) gut gelungene Darstellung des *paraverbal code* nach Chaume (2004) in „Behind Her Eyes“ findet sich hingegen in Folge 6 (08:05-09:15), als Rob beim gemeinsamen Kochen mit Adele und David einen französischen Akzent nachahmt:

Nummer	Sichtbare Handlung	Transkription Musik, Geräusche
6	Rob, Adele und David sind gemeinsam in der Küche. Rob erklärt David und Adele, wie er das Essen vorbereitet und kocht, und setzt dabei einen übertriebenen französischen Akzent auf.	Hintergrundkulisse der Küche (Klirren des Bestecks, der Töpfe etc.)

Nummer	Transkription Dialog		
6	Englisch	Deutsch	Französisch
	[mit französischem Akzent] The trick with the potatoes is to drop them in the boiling water, and then you can hear the potatoes scream!	[mit französischem Akzent] Nun, der Trick mit den Kartoffeln ist, äh, man werfe sie in das kochende Wasser, und dann hört man die Kartoffeln schreien!	Au menu ce soir, les pommes sautées du chef. Je m'appête à plonger les patates dans l'eau bouillante. Attention, vous allez les entendre hurler de douleur !

Nummer	Transkription SDH		
6	Englisch	Deutsch	Französisch
	[in French accent] The trick with the potatoes	[französischer Akzent] Nun, der Trick mit den Kartoffeln ist, äh,	Au menu ce soir, les pommes sautées du chef.
	is to drop them in the boiling water,	man werfe sie in das kochende Wasser,	Je m'appête à plonger les patates dans l'eau bouillante.
	and then you can hear the potatoes scream!	und dann hört man die Kartoffeln schreien!	Attention, vous allez les entendre hurler de douleur !

Die deutsche Synchronisierung behält in dieser Szene die gleiche Strategie bei wie die englischsprachige Originaltonspur, nämlich den französischen Akzent als humoristischen Marker für Robs Charakter zu verwenden. Da Frankreich gemeinhin als das Land der Haute Cuisine, also der hohen Kochkunst, bekannt ist, spielt Rob also darauf an, dass die Mahlzeit, die er gerade vorbereitet, besonders exquisit ist. Im weiteren Verlauf der Szene verwenden sowohl Adele als auch David französische Wörter und Phrasen, zum Beispiel „Monsieur Chef“ oder „Et voilà“, als sie mit Rob sprechen, und tragen so auch zur humoristischen Äußerung bei. Die Übermittlung von Humor vollzieht sich in dieser Szene demnach nicht nur durch den Akzent – die generelle gute Laune und der Spaß der Charaktere sind eindeutig am Bildschirm zu sehen, weshalb der Akzent hier vermutlich nur eine unterstützende Rolle spielt.

In der französischen Sprachfassung wurde schon bei der Synchronisierung entschieden, keine besonderen sprachlichen Merkmale einzubringen – Rob spricht also für die französischsprachigen Zuseher:innen so wie auch sonst in der Serie. Zu einem späteren Zeitpunkt in dieser Szene wird erwähnt, dass er übertrieben bzw. gestellt ([exagérant]) spricht, aber nicht mit einem besonderen Akzent. Diese Lösung finde ich gut gelungen, da die humoristische Äußerung in dieser Situation nicht primär am Akzent festgemacht wird, sondern an der generellen Interaktion der drei Charaktere. Die Hauptfunktion der Szene – also die Übertragung von Humor – bleibt erhalten, auch wenn Rob in der französischen Sprachversion

„nur“ gestellt spricht, anstatt einen bestimmten Akzent aufzulegen. Es sollte also in allen drei Sprachfassungen für die hörgeschädigten Rezipient:innen verständlich und nachvollziehbar sein, wieso die Charaktere bei der Zubereitung des Essens so viel Spaß haben – die *narrative clarity* (vgl. Vandaele 2018) sollte auch ohne die Akzentbeschreibung gegeben sein.

5.3.2 Sprecher:innen-Identifikation

Die Sprecher:innen-Identifikation zählt zu den Hauptmerkmalen von SDH-Untertiteln. Sie wird vor allem dann wichtig, wenn die Sprecher:innen gerade nicht im Bild zu sehen sind – laut Netflix-Richtlinien soll sie auch nur dann erfolgen. Im nachfolgenden Beispiel ist ein deutlicher Unterschied zwischen den einzelnen Sprachfassungen zu erkennen, welche auch Auswirkungen auf die filmische Narration haben kann. Da in diesem Beispiel speziell die Verwendung von Speaker-IDs veranschaulicht werden soll, habe ich aus Platzgründen einige Stellen ausgelassen – die hier genannten Äußerungen und Untertitel folgen also nicht direkt aufeinander.

Nummer	Sichtbare Handlung	Transkription Musik, Geräusche
7	Louise macht sich im Schlafzimmer fertig; sie tanzt dabei und betrachtet sich im Spiegel. Adam kommt ins Zimmer, und sie unterhalten sich. [...] Laura (Louises und Adams Nachbarin) kommt vorbei, um auf Adam aufzupassen, während Louise ausgeht.	Lebhafte, fröhliche Musik („Never Forget You“ von Noisettes)

Nummer	Transkription Dialog		
7	Englisch	Deutsch	Französisch
	[Adam] What are you doing?	[Adam] Was machst du da?	[Adam] Qu'est-ce que tu fais ?
	[...]	[...]	[...]
	[Louise] Yeah, good. Thanks.	[Louise] Ja, alles klar. Danke.	[Louise] Oui, ça va. Merci.
	[...]	[...]	[...]
	[Laura] How you doing, A?	-[Laura] Wie geht's, Adam?	[Laura] Ça va, Adam ?
	[Adam] Good.	-[Adam] Gut.	[Adam] Ça va.

Nummer	Transkription SDH		
7	Englisch	Deutsch	Französisch
	[Adam] What are you doing?	[Junge] Was machst du da?	[garçon] Qu'est-ce que tu fais ?

	[Louise] Yeah, good. Thanks.	[Mutter] Ja, alles klar. Danke.	[femme] Oui, ça va. Merci.
	-[Laura] How you doing, A?	-[Frau] Wie geht's, Adam?	- [femme 2] Ça va, Adam ?
	-[Adam] Good.	-[Adam] Gut.	- Ça va.

In dieser Szene wird die Wichtigkeit des *syntactic code* für die SDH-Untertitelung veranschaulicht. Da sie gleich zu Beginn der ersten Folge auftritt (0:37-1:07), kennen die Zuseher:innen die Namen der Charaktere noch nicht. Laut Netflix-Richtlinien (und laut gängigen Konventionen) sollen die Untertitel im Idealfall das gleiche Maß an Informationen zur Verfügung stellen wie die Audiospur, und die deutsch- und französischsprachige Version tun genau das. Am Anfang, bevor jegliche Namen genannt werden, werden die Charaktere als „Junge“, „Mutter“, „femme“, „femme 2“ et cetera bezeichnet. Nachdem die Namen in der Audiospur zum ersten Mal genannt wurden, werden auch die Speaker-IDs in den Untertiteln angepasst und bezeichnen dann die Charaktere auch mit ihren richtigen Namen (z.B. in Untertitel 3 im oben genannten Beispiel).

In den englischsprachigen Untertiteln werden die Charaktere sofort mit ihren Namen benannt, was den Richtlinien und Konventionen widerspricht. Den SDH-Rezipient:innen stehen also mehr Informationen zur Verfügung, als der Ausgangstext bekanntgibt – und damit auch mehr Informationen, als sie eigentlich bekanntgeben sollen. Das eigenständige Suchen nach relevanter Information ist eines der Hauptmerkmale von Narration in multimodalen Produkten und soll bei den Zuseher:innen Neugierde und Spannung auslösen (vgl. Vandaele 2018). Während das hörende und nicht-englischsprachige gehörlose Publikum also auf diese Informationen wartet und sich selbstständig auf die Suche nach diesen Informationen macht, um mehr über die Charaktere zu erfahren, werden den englischsprachigen gehörlosen Zuseher:innen die Namen der Charaktere sofort genannt. Ihre Sicht auf die Narration der Serie ist eine andere als die der anderen Rezipient:innen, denen die Information nicht zur Verfügung steht. Die Sprachfassungen weisen also unterschiedliche Grade an *temporary opacity* (vgl. Vandaele 2018) auf, was wiederum dazu führen kann, dass in den Zuseher:innen ein unterschiedliches Maß an Neugierde ausgelöst wird.

5.3.3 Andere paraverbale Merkmale

Ein Beispiel für die Beschreibung anderer paraverbaler Merkmale findet sich in Folge 3 (1:44-3:09), gleich im Anschluss an die Szene, die in Punkt 5.2 analysiert wurde. Auch hier vollzieht sich ein großer Teil der Kommunikation über Geräusche und paraverbale Merkmale, da es

ansonsten in dieser Szene nicht viel Dialog gibt. (Aus Platzgründen werde ich auch hier nur die paraverbalen Merkmale transkribieren; die gesamte Transkription findet sich in Punkt 5.2.)

Nummer	Sichtbare Handlung	Transkription Musik, Geräusche
8	Szene 1: David verlässt früh morgens das Haus und versucht dabei leise zu sein, um Adele nicht aufzuwecken. Szene 2: Die Szene wechselt ins Büro, wo Louise die Kaffeemaschine bedient und sich anschließend in Davids Büro setzt. Als sie aufsteht, betritt David plötzlich sein Büro.	Szene 1: Anfangs keine Musik, Schritte, Stufe knarzt, Schlüssel klimpern, Adele atmet laut, Türe öffnet sich und fällt zu. Melancholische Klaviermusik beginnt, Vogelzwitschern, Hundebellen. Szene 2: Melancholische Musik spielt weiter, Klappern der Kaffeemaschine, Louise seufzt.

Nummer	Transkription Dialog		
8	Englisch	Deutsch	Französisch
	Beide stammeln.	Beide stammeln.	Beide stammeln.
	[David] I'll just ... Sorry.	[David] Entschuldigung.	[David] Désolé.
	[Louise] Sorry.		[Louise] Pardon.

Nummer	Transkription SDH		
8	Englisch	Deutsch	Französisch
	[exhales]	[seufzt]	[souponir]
	[Louise exhales]	-[Louise prustet]	[Louise souffle]
	[both gasp]	-[holt tief Luft]	
	-Oh.	Oh... Ähm...	[sursauts]
	-Oh.		
	Eh... eh...		[balbutiements]
	-I'll just... Sorry. [chuckles]	-[Louise] Äh...	- Désolé.
	-Sorry. [chuckles]	-Entschuldigung.	- Pardon.

Auch hier sind verschiedene Darstellungen des *paralinguistic code* nach Chaume (2004) sichtbar. Zum einen werden die diegetischen paraverbalen Äußerungen der Charaktere – das Ausatmen, das Prusten und das Luftholen – in allen drei Sprachfassungen in Klammern angegeben und begleiten die sichtbaren Handlungen im Bild. Auf diese Weise können auch hörgeschädigte Zuseher:innen die paraverbalen Äußerungen nachvollziehen; die diegetischen

multimodalen Bestandteile der Szene werden in ihrer gesamten Ausprägung in den Untertitel reproduziert.

An den Untertiteln in dieser Szene ist vor allem das Ende sehr interessant und relevant für die Analyse. Während in den deutschen und englischen Untertiteln das Stammeln von Louise und David am Ende in Worten ausgeschrieben wird („Oh“, „Ähm“, „Eh ...“), sind diese Teile in der französischsprachigen Version als paraverbale Merkmale beschrieben, also [sursauts] (Aufschrecken) und [balbutiements] (Stammeln). Meiner Meinung nach beschreiben beide Arten der Transkription die Szene gut – die beiden führen nicht wirklich ein Gespräch, sondern reden mehr oder weniger aneinander vorbei, weshalb der genaue wörtliche Inhalt der Äußerungen in diesem Fall nicht das wichtigste Element ist. Anhand beider Schreibweisen lässt sich gut die Überraschung und Überwältigung der Charaktere ablesen, egal, ob sie nun als Dialog zwischen den beiden dargestellt ist oder durch paraverbale Situations- und Kontextangaben. Die *narrative clarity* nach Vandaele (2018) zielt in diesem Fall nicht darauf ab, den Dialog genau wiederzugeben, sondern die allgemeine Verwirrung der Charaktere neben der bildlichen Darstellung auch schriftlich zu reproduzieren; dies gelingt durch beide Arten der Transkription. Die französischsprachigen Untertitel sind in diesem Beispiel um einiges kürzer als die deutsch- und englischsprachigen, beinhalten dabei jedoch das gleiche Maß an Information – das könnte auch bedeuten, dass die Zuseher:innen der französischen Sprachfassung weniger mentalen Aufwand bei der Rezeption der Untertitel aufbringen müssen.

Im Moment gibt es meines Wissens weder in der Forschung noch in der Praxis eindeutigen Konsens darüber, welche Darstellungsweise von Menschen mit Hörbehinderung bevorzugt wird. Wenn dies wirklich der Fall ist und beide Schreibweisen gleich gut verständlich sind, handelt es sich hierbei wahrscheinlich um eine Frage der persönlichen Präferenz oder um Konventionen der einzelnen Sprachregionen.

5.4 Diskussion der Ergebnisse

Bevor ich noch einmal auf die oben beschriebenen Beispiele und ihre Auswirkungen auf die filmische Narration eingehe, möchte ich an dieser Stelle einige generelle Anmerkungen zu den verschiedenen Sprachfassungen der SDH-Untertitel machen. Als erstes möchte ich auf die sprachliche Bildhaftigkeit eingehen; mir schienen die englisch- und deutschsprachigen Untertitel um einiges bildhafter als die französischsprachigen, da sie genauere Beschreibungen und vielfältigeres Vokabular enthielten. Ein Beispiel hierfür ist die in Kapitel 5.2 beschriebene

Szene, welche in der deutschen und englischen Sprachfassung um einiges detaillierter in den Untertiteln reproduziert wurde.

Eine weitere Tatsache, die mir bei der Analyse ins Auge gestochen ist, war die variierende Anzahl an Situations- und Kontextangaben zwischen den einzelnen Sprachversionen. Im Vergleich zu den anderen zwei Sprachfassungen standen den französischsprachigen Zuseher:innen in jeder Folge deutlich weniger Beschreibungen zur Verfügung. Während die durchschnittliche Anzahl an Situations- und Kontextangaben der deutschen und englischen Sprachfassungen sehr ähnlich ist (251 bzw. 256 Beschreibungen pro Folge), liegt der Durchschnitt in der französischen Sprachfassung bei 167 Beschreibungen pro Folge und ist damit deutlich geringer. Das muss nicht zwangsläufig bedeuten, dass die französischsprachigen Rezipient:innen die Narration der Serie schlechter nachvollziehen können als die anderen – wie bereits erwähnt, können gehörlose Zuseher:innen Kontextinformationen, die in der Untertitelung fehlen, anhand von eigenen Erfahrungen und Weltwissen erschließen (vgl. Remael & Reviers 2018; Tiittula & Hirvonen 2019). Dennoch ist es möglich, dass das englisch- und deutschsprachige Publikum weniger mentale Anstrengung bei der Rezeption der Serie aufwenden muss, da ihnen mehr Informationen in Hinblick auf Situation und Kontext zur Verfügung stehen.

In Bezug auf Musikbeschreibungen ließen sich in den drei Sprachregionen unterschiedliche Strategien erkennen. Während die Titelmelodie nur im Deutschen näher beschrieben wurde, erwähnten die anderen beiden Sprachfassungen sie nicht. Auch die generischen Musikbeschreibungen (für Hintergrundmusik) sind in der deutschen Version am ausführlichsten; hier wurden auch Veränderungen der Musik in die Untertitelung miteinbezogen, etwa, wenn die Hintergrundmusik lauter oder leiser wurde oder sich die Stimmung veränderte (z.B. [Musik wird düsterer]). Die englisch- und französischsprachigen Untertitel erwähnten in diesem konkreten Fall nur die generelle Stimmung der Musik, es erfolgte aber keine Beschreibung etwaiger Veränderungen. Das bedeutet in weiterer Folge, dass für deutschsprachige hörgeschädigte Zuseher:innen die Stimmung, die sowohl durch die Titel- als auch durch die Hintergrundmusik vermittelt werden soll, wahrscheinlich besser bzw. mit weniger mentaler Anstrengung nachvollziehbar ist, da ihnen mehr Informationen zur Verfügung stehen.

Bei Musiktiteln und den dazugehörigen Texten wiederum zeigte sich eine andere Strategie. Laut Netflix-Richtlinien sollen nur jene Lieder untertitelt werden, die in der gleichen Sprache verfügbar sind wie die Untertitelung. Da „Behind Her Eyes“ eine englischsprachige Produktion ist, wurden auch nur englische Songs in der Serie verwendet. Das bedeutet also,

dass nur die englischsprachigen Gehörlosen-Untertitel die Songtexte auch reproduzieren, während die deutsch- und französischsprachigen Untertitel das nicht tun. In der vorliegenden Analyse wurde ein Spezialfall untersucht, in dem die deutschen Untertitel eine interlinguale Übersetzung des englischen Liedtextes bereithielten, sodass auch deutschsprachige Zuseher:innen auf die Information, die durch den Song vermittelt werden soll, zugreifen können.

Da Musik ein wichtiger Bestandteil der filmischen Narration ist (vgl. Garwood 2006), ist es unabdingbar, sie auch in den SDH-Untertiteln zu reproduzieren. Das gilt sowohl für Szenen, in denen sich die Narration hauptsächlich durch Musik vollzieht – zum Beispiel, wenn ein Liedtext deutlich zu hören ist –, als auch für Szenen, in denen die Musik als unterstützender Hintergrundfaktor für die Narration dient. Sie sollte daher jedenfalls inkludiert werden, damit hörgeschädigte Zuseher:innen das multimodale Produkt in gleichem (bzw. in vergleichbarem) Maße nachvollziehen können wie hörende Personen. Die Beschreibung von Musik könnte auch Situationen verhindern, in denen das Publikum nicht versteht, wieso ein bestimmtes Lied in einem bestimmten Kontext verwendet wurde (vgl. Desilla 2009). Außerdem kann das Nicht-Übersetzen bzw. die Nicht-Transkription von Musik dazu führen, dass multimodale Texte nicht oder nur oberflächlich verstanden werden (vgl. Di Giovanni 2008), was in jedem Fall vermieden werden sollte.

In Bezug auf Geräuschbeschreibungen ließen sich in den drei Sprachfassungen ebenfalls deutliche Unterschiede feststellen. In diesem Fall enthielten die deutsch- und englischsprachigen Untertitel die meisten Informationen; die narrative Welt, die durch die französischsprachigen Untertitel dargestellt wurde, war hingegen weniger lebendig, da zum Beispiel nicht-diegetische Hintergrundgeräusche ausgelassen wurden. Die diegetischen Geräusche wurden in allen drei Sprachfassungen beschrieben – hörgeschädigte Zuseher:innen sollten also jedenfalls der wesentlichen Narration der Serie folgen können. Das Nicht-Beschreiben von auditiv vermittelten Informationen kann jedoch dazu führen, dass die hörgeschädigten Zuseher:innen ein anderes Bild von der Welt der Charaktere haben als die hörenden, welche sehr wohl Zugriff auf die Hintergrundgeräusche haben. Denn auch wenn zum Beispiel ein Hundebellen keine direkte diegetische Funktion in einer Szene hat, so macht sie diese doch lebensechter. Hörgeschädigte Zuseher:innen könnten demnach denken, dass in dieser Szene keine anderen Geräusche vorhanden sind, was wiederum zu größerer mentaler Anstrengung führen könnte, wenn sie sich die narrative Welt selbst erschließen müssen. Es wäre auch denkbar, dass sie die Narration nicht in gleichem Maße nachvollziehen können wie die hörenden Zuseher:innen, was dem Sinn der Barrierefreiheit entgegenwirken würde.

Auch die Beschreibungen der Charaktere wiesen einige Unterschiede auf. Ein Beispiel, in dem ein paraverbales Merkmal von Davids Charakter – sein Akzent – zwar zu einem zentralen Bestandteil des Dialogs wurde, aber in keiner Sprachfassung in den Untertiteln erwähnt wurde, zeigte, dass eine solche Auslassung potenziell zu einem höheren mentalen Aufwand für die hörgeschädigten Zuseher:innen führen könnte. Es wäre auch denkbar, dass das Publikum diese Aussage nicht nachvollziehen kann und sie daher als verwirrend empfindet, was ebenfalls (vermeidbare) höhere mentale Anstrengung mit sich tragen würde.

In einem anderen Beispiel wurde ein von Rob aufgesetzter Akzent beschrieben, der als narrativer Marker für Humor in der Szene verwendet wurde. Die englisch- und deutschsprachigen Untertitel erwähnten diesen in eckigen Klammern ([französischer Akzent]), wodurch die Assoziation auch für gehörlose Personen zugänglich war, welche die humoristische Aussage dadurch verstehen konnten. Im Gegensatz dazu wurde in der französischen Version der Untertitel nur erwähnt, dass Rob übertrieben bzw. gestellt ([exagérant]) spricht. Da während dieser Aussage die Hauptcharaktere auch am Bildschirm zu sehen sind, unterstützt die sichtbare Handlung die humoristische Äußerung. Es sollte daher meiner Meinung nach in allen drei Sprachfassungen nachvollziehbar sein, wieso die Hauptfiguren sich in dieser Szene so sehr amüsieren.

Andere paraverbale Merkmale wurden im ausgewählten Beispiel zwar auf unterschiedliche Weise transkribiert – durch Worte (englisch und deutsch) oder durch Beschreibungen in Klammern (französisch) –, jedoch machten beide Arten der Beschreibung meiner Meinung nach die Situation nachvollziehbar. Da es in dieser Szene nicht hauptsächlich um den Dialog ging, sondern mehr um das Darstellen der vorherrschenden Stimmung, ist es wohl eher eine Frage der persönlichen Präferenz bzw. von sprachraumabhängigen Konventionen, welche Darstellungsweise gewählt werden sollte.

Im letzten analysierten Beispiel aus „Behind Her Eyes“ ging es um Speaker-IDs, also das Erwähnen der Sprecher:innen in den Situations- und Kontextangaben. Laut Netflix-Richtlinien und Konventionen sollen die Untertitel im besten Fall immer das gleiche Maß an Information bereitstellen wie die Originaltonspur. Das bezieht sich auch auf Namen – solange die Namen der Charaktere nicht explizit genannt werden, sollten sie auch in den Untertiteln nicht vorkommen. Die deutsch- und französischsprachigen Untertitel hielten sich an diese Vorgaben und nannten die Figuren zu Beginn nicht mit Namen; die englischsprachigen Untertitel hingegen enthielten von Beginn an die Namen der Hauptcharaktere, was den Richtlinien und Konventionen widerspricht. Das Vorenthalten von Information kann ein wichtiger Bestandteil der filmischen Narration sein, da es in den Zuseher:innen Neugierde

auslösen soll. In diesem Aspekt unterscheiden sich die Sprachfassungen also, was auch dazu führen könnte, dass die Rezipient:innen jeweils einen anderen Eindruck von der narrativen Welt der Charaktere haben.

Zusammenfassend lässt sich sagen, dass in den meisten Fällen die hörgeschädigten Zuseher:innen die filmische Narration in „Behind Her Eyes“ nicht in gleichem Maße nachvollziehen konnten wie hörende Personen. Auch hier steht also, wie in so vielen Fällen, das audiovisuelle Produkt behinderten Personen nicht uneingeschränkt zur Verfügung, während able-bodied Menschen ohne Weiteres darauf zugreifen können. Selbstverständlich sind einige Bausteine eines multimodalen Textes schwieriger schriftlich darzustellen als andere – Dialekte und Geräusche sind zum Beispiel oft schwieriger zu beschreiben als ein Dialog –, jedoch ist es für die Narration eines Textes in den meisten Fällen förderlich, zumindest eine Beschreibung einzubauen. Denn auch wenn Rezipient:innen zum Beispiel den Text eines Liedes nicht verstehen können, so können sie durch eine Musikbeschreibung zumindest die Stimmung des Liedes nachvollziehen, was wiederum für das Verständnis einer gezeigten Szene hilfreich sein kann. Für die Situations- und Kontextangaben in „Behind Her Eyes“ waren vorwiegend jene Codes nach Chaume (2004) ausschlaggebend, die zu Beginn der Analyse als relevant identifiziert wurden: der *paralinguistic code* und der *musical and special effects code*. Es stellte sich aber heraus, dass auch der *syntactic code* und der *sound arrangement code* großen Einfluss auf die Untertitel haben können. In allen Beispielen waren selbstverständlich auch andere Codes vorhanden, etwa der *linguistic code*, doch diese spielten im Rahmen der durchgeführten Analyse eher eine unterstützende Rolle.

6. Conclusio

Audiovisuelle Produkte wie Filme oder Serien zeichnen sich dadurch aus, dass aus dem Zusammenspiel mehrerer bedeutungsgebender Codes (Bild, Ton, Musik, Sprache) ein zusammenhängendes Gesamtbild entsteht, welches mehr als nur die Summe seiner Teile ist. Aus diesem Grund zählen audiovisuelle Texte zur multimodalen Kommunikation, welche auf dem gleichzeitigen Auftreten verschiedener Modi basiert und daher den geeigneten theoretischen Rahmen für die vorliegende Masterarbeit darstellt. Das Konzept der Multimodalität wurde lange nicht in die Translationswissenschaft integriert und hat sich dennoch innerhalb der letzten Jahrzehnte zu einem der Hauptforschungsinteressen entwickelt; zahlreiche Wissenschaftler:innen beschäftigen sich heute etwa mit SDH-Untertitelung und Audiodeskription oder betreiben Forschung zu Analysemodellen für multimodale Texte (unter anderem multimodale Korpora und multimodale Transkription).

Die SDH-Untertitelung ist außerdem ein Teilbereich der barrierefreien Kommunikation, da sie den hörbaren sprachlichen Code in audiovisuellen Produkten schriftlich reproduziert und ihn so hörgeschädigten und gehörlosen Menschen zugänglich macht. Sie ist also ein wichtiges Instrument für die Teilhabe und Partizipation am alltäglichen Leben. Die Zielgruppe dieser Form der Untertitelung ist besonders heterogen – sie umfasst nicht nur die eben genannten Menschen mit Hörbehinderung, sondern zum Beispiel auch Menschen mit Lernschwierigkeiten oder Sprachenlerner:innen. Darin besteht auch eine der größten Herausforderungen dieser Form der Untertitelung: Sie kann nicht gleichzeitig die Bedürfnisse aller Zielgruppen in gleichem Maße befriedigen.

Die Translation audiovisueller Produkte ist heute einer der Schwerpunkte der Translationswissenschaft; auch sie wurde lange nicht wissenschaftlich diskutiert, da sie vom Feld der traditionellen Übersetzung abweicht. Sie kann unterschiedliche Ausprägungen annehmen, zu denen auch die SDH-Untertitelung zählt. Als intralinguale Übersetzung der Tonspur inkludiert diese neben dem gesprochenen Dialog auch Musik- und Geräuschbeschreibungen sowie Angaben zu Sprecher:innen (paraverbale Merkmale, Speaker-IDs). Diese Form der Untertitelung stellt ein Hauptforschungsinteresse der *Media Accessibility* dar – unter anderem wird untersucht, wie Kinder SDH rezipieren, inwiefern Symbole für die Verständlichkeit hilfreich sein können oder wie sich die filmische Narration in audiovisuellen Produkten mit Untertiteln vollzieht. Auch Automatisierungsprozesse sind von großem wissenschaftlichen und wirtschaftlichen Interesse. Viele Forscher:innen kritisieren jedoch, dass in diesem Bereich bis heute keine international gültigen Normen existieren, weshalb sowohl

sprachraumübergreifend als auch innerhalb der einzelnen Sprachregionen große Unterschiede in den Untertitelungen bestehen können.

Mit ebendiesen Unterschieden bezüglich SDH-Untertitelung beschäftigte sich die vorliegende Masterarbeit. Genauer hatte sie das Ziel, herauszufinden, inwiefern die Situations- und Kontextangaben Einfluss auf die filmische Narration in der Serie „Behind Her Eyes“ nehmen. Zu diesem Zweck wurde die Serie zunächst multimodal transkribiert, bevor sie anhand des Analysemodells für audiovisuelle Texte nach Chaume (2004) untersucht wurde. Dieses Modell vereint translations- und filmwissenschaftliche Aspekte und eignet sich daher besonders gut für die Untersuchung der Untertitelung.

Die Analyse wurde in drei Kapitel unterteilt: Musik, Geräusche und Charaktere. In Hinblick auf alle drei Aspekte zeigte sich im Rahmen der Untersuchung sowohl, dass deutschsprachige Zuseher:innen im Durchschnitt die höchste Anzahl an Untertiteln pro Folge zur Verfügung hatten, als auch, dass diese Untertitel ausführlichere Informationen enthielten als die englisch- und französischsprachigen. Ein Beispiel hierfür sind etwa die generischen Musikbeschreibungen, welche in der deutschen Sprachfassung genauer und öfter beschrieben wurden (z.B. [Musik wird düster]) als in der englisch- und französischsprachigen ([melancholy music playing] bzw. [musique mélancolique]). Im Bereich der Geräuschbeschreibungen waren die deutsch- und englischsprachigen Untertitel deutlich informationsreicher als die französischsprachigen, da in letzterer Sprachfassung oftmals nicht-diegetische Hintergrundgeräusche wie Hundebellen ausgelassen wurde. Diese Geräusche sind zwar nicht direkt handlungsleitend, machen die narrative Welt der Charaktere jedoch lebensechter und sind daher trotzdem ein wichtiger Bestandteil des audiovisuellen Produktes. Hörgeschädigte Zuseher:innen, die nicht auf umfangreiche Untertitel zurückgreifen können, könnten demnach entweder denken, dass in einer Szene keine Geräusche vorhanden sind, oder sie müssen die Narration selbst erschließen, was eine höhere mentale Anstrengung zur Folge hat.

Laut Richtlinien und gängigen Konventionen in der Praxis sollen SDH-Untertitel nur so viele Informationen enthalten, wie die Tonspur bereithält – das gilt auch für die Namen der Charaktere. In der analysierten Serie wurden die Namen der Hauptcharaktere in der englischsprachigen Version der Untertitel jedoch sofort bekanntgegeben, während sie in den anderen zwei Sprachen zu Beginn geheim blieben und erst dann erwähnt wurden, wenn die Tonspur sie offenlegte. Das Vorenthalten von Information kann ein wichtiges Instrument für die Narration sein, da es in den Zuseher:innen Neugierde auslösen soll – in diesem Beispiel könnten die genannten Unterschiede also dazu führen, dass die Rezipient:innen jeweils einen anderen Blick auf die narrative Welt der Charaktere haben.

Die Untersuchung im Rahmen der Masterarbeit zeigte, dass den hörgeschädigten Rezipient:innen größtenteils weniger Situations- und Kontextinformationen zur filmischen Narration von „Behind Her Eyes“ zur Verfügung standen als hörenden Personen. Die beschriebenen Ergebnisse müssen jedoch nicht zwangsläufig bedeuten, dass hörgeschädigte Rezipient:innen die narrative Welt auch wirklich schlechter verstehen als hörende. Wie bereits erwähnt, wurde in zahlreichen Studien festgestellt, dass Menschen mit Hörbehinderung durch eigenes Weltwissen und Erfahrungen viele Informationen selbst erschließen können, ohne dass diese in den Untertiteln reproduziert wurden (vgl. Remael & Reviers 2018; Tiittula & Hirvonen 2019). Um festzustellen, welchen Einfluss die Situations- und Kontextangaben in Untertiteln aus der Sicht von hörgeschädigten Personen haben, müsste also eine Rezeptionsstudie durchgeführt werden. Das würde auch sicherstellen, dass – wie so oft bei barrierefreien Angeboten – nicht nur über die Zielgruppe gesprochen wird, sondern ein aktiver Diskurs mit der Zielgruppe über Barrierefreiheit herrscht. Sie sollte aktiv in den Untertitelungsprozess eingebunden werden, wie das bei der Audiodeskription bereits heute der Fall ist, sodass die Untertitel besser an ihre Bedürfnisse angepasst werden können.

Ein weiterer wichtiger Punkt wäre die Integration von Forschungsergebnissen in die Praxis, denn obwohl schon zahlreiche Studien die Möglichkeiten und Grenzen intralingualer Untertitel aufgezeigt haben, finden die Erkenntnisse nur selten Eingang in die Praxis. Die Gründe dafür können vielzählig sein; es ist allerdings sehr wahrscheinlich, dass begrenzte zeitliche und finanzielle Ressourcen dafür ausschlaggebend sind. Der Untertitelungsprozess findet klassischerweise in der Postproduktion statt – den Untertitler:innen stehen also nur die vollendeten Filme zur Verfügung, die gegebenenfalls bereits synchronisiert wurden, ohne dass auf die Untertitelbarkeit Rücksicht genommen wurde (z.B. zu hohe Sprechgeschwindigkeit, Verwendung von Akzenten). Das kann dazu führen, dass Informationen verloren gehen, weil sie aufgrund von Platzmangel oder zeitlichen Einschränkungen nicht in die Untertitel inkludiert werden können, was wiederum zur Folge haben kann, dass hörgeschädigte Personen die Handlung nicht in gleichem Maße nachvollziehen können wie das hörende Publikum.

Es ist also noch ein langer Weg, bis man wirklich von Media Accessibility für alle beteiligten Zielgruppen sprechen kann. Daher wäre es wünschenswert, dass vorliegende Forschungsergebnisse zu SDH und anderen barrierefreien Services in Zukunft auch in der Praxis angewendet werden, sodass die primären Zielgruppen der barrierefreien Angebote – in diesem Fall Personen mit Hörbehinderung – auch von ihnen profitieren können. In weiterer Folge könnten dann auch die unzähligen weiteren Zielgruppen, also zum Beispiel Menschen mit Lernschwäche oder Sprachenlerner:innen, einen Nutzen daraus ziehen.

Primärquellen

- Richter Strand, Erik (2021). Behind Her Eyes: Chance Encounters (Staffel 1, Folge 1) [Folge einer TV-Serie, 50 min.]. Vereinigtes Königreich: Netflix.
- Richter Strand, Erik (2021). Behind Her Eyes: Lucid Dreaming (Staffel 1, Folge 2) [Folge einer TV-Serie, 50 min.]. Vereinigtes Königreich: Netflix.
- Richter Strand, Erik (2021). Behind Her Eyes: The First Door (Staffel 1, Folge 3) [Folge einer TV-Serie, 51 min.]. Vereinigtes Königreich: Netflix.
- Richter Strand, Erik (2021). Behind Her Eyes: Rob (Staffel 1, Folge 4) [Folge einer TV-Serie, 54 min.]. Vereinigtes Königreich: Netflix.
- Richter Strand, Erik (2021). Behind Her Eyes: The Second Door (Staffel 1, Folge 5) [Folge einer TV-Serie, 47 min.]. Vereinigtes Königreich: Netflix.
- Richter Strand, Erik (2021). Behind Her Eyes: Behind Her Eyes (Staffel 1, Folge 6) [Folge einer TV-Serie, 51 min.]. Vereinigtes Königreich: Netflix.

Bibliografie

- Azenha, Joao & Moreira, Marcelo (2012). Translation and Rewriting: Don't Translators „adapt“ when they „translate“? In: Raw, Laurence (Hg.) *Translation, adaptation and transformation*. London: Bloomsbury , 61-80.
- Baldry, Anthony & Thibault, Paul (2006). *Multimodal Transcription and Text Analysis*. London: Equinox.
- BBC (2021). Subtitle Guidelines. <https://bbc.github.io/subtitle-guidelines/#Colours> (Stand: 18.11.2021)
- Benner, Uta & Herrmann, Annika (2018). Gebärdensprachdolmetschen. In: Maaß, Christiane & Rink, Isabel (Hg.) *Handbuch Barrierefreie Kommunikation*. Berlin: Frank & Timme, 381-402.
- Berthele, Raphael (2000). Translating African-American Vernacular English into German: The problem of ‚Jim‘ in Mark Twain's Huckleberry Finn. *Journal of Sociolinguistics* 4 (4), 588-613.
- Besthelp (o.J.). Lexikon – Sehbehinderung. <https://www.besthelp.at/lexikon/sehbehinderung> (Stand: 23.02.2022)

- Boria, Monica & Tomalin, Marcus (2020). Introduction. In: Boria, Monica; Carreres, Ángeles; Noriega-Sánchez, María & Tomalin, Marcus (Hg.) *Translation and multimodality: beyond words*. London: Routledge, 1-23.
- Bosseaux, Charlotte (2008). Buffy the Vampire Slayer: Characterization in the Musical Episode of the TV Series, *The Translator* 14 (2), 343-372.
- BSVÖ (o.J.). Statistische Daten. <https://www.blindenverband.at/de/information/augengesundheit/97/Statistische-Daten> (Stand: 23.02.2022)
- Bundesgesetz über die Gleichstellung von Menschen mit Behinderungen (Bundes-Behindertengleichstellungsgesetz – BGStG), BGBl. I Nr. 82/2005.
- Chaume, Frederic (2004). Film Studies and Translation Studies: Two Disciplines at Stake in Audiovisual Translation. *Meta: Translator's Journal* 49 (1), 12-24.
- Civera, Clara (2005). Introducing emoticons and pictograms in SDHH. *Media for All conference*.
- Civera, Clara & Orero, Pilar (2010). Introducing icons in subtitles for the deaf and hard of hearing: Optimising reception? In: Matamala, Anna & Orero, Pilar (Hg.) *Listening to Subtitles: Subtitles for the Deaf and Hard of Hearing*. Bern: Peter Lang, 149-162.
- Conseil Supérieur de l'Audiovisuel (2012). Dossier sous-titrage: La forme, le contenu et la technique. <http://www.sirtin.fr/sirtin/wp-content/uploads/soustitrage.pdf> (Stand: 18.11.2021)
- Danan, Martine (1992). Reversed Subtitling and Dual-Coding Theory: New Directions for Foreign Language Instruction. *Language Learning* 42(4), 497–527.
- Das Erste (2021). Untertitel-Standards von ARD, ORF, SRF, ZDF. <https://www.daserste.de/specials/service/untertitel-standards100.html> (Stand: 18.11.2021)
- Desilla, Louisa (2009). *Towards a Methodology for the Study of Implicatures in Subtitled Films: Multimodal Construal and Reception of Pragmatic Meaning Across Cultures*. Dissertation, University of Manchester.
- Díaz Cintas, Jorge (1998). La labor subtituladora en tanto que instancia de traducción subordinada. In: Pilar Orero (Hg.) *Actes del III Congr s Internacional sobre Traducci *. Barcelona: Universitat Aut noma de Barcelona, 83–90.
- D az Cintas, Jorge & Remael, Aline (2007). *Audiovisual Translation: Subtitling*. London: Routledge.

- Díaz Cintas, Jorge; Orero, Pilar & Remael, Aline (2007). Media for all: a Global Challenge. In: Díaz Cintas, Jorge; Orero, Pilar & Remael, Aline (Hg.) *Media for All: Subtitling for the Deaf, Audio Description, and Sign Language*. Amsterdam: Rodopi, 11-20.
- Díaz Cintas, Jorge & Remael, Aline (2021). *Subtitling: Concepts and Practices*. London: Routledge. <https://doi-org.uaccess.univie.ac.at/10.4324/9781315674278> (Stand: 27.09.2021)
- Di Giovanni, Elena (2008). The American Film Musical in Italy: Translation and Non-translation. *The Translator* 14 (2), 295-318.
- Ein Blog von Vielen (2020). *Das Gegenteil von "behindert"*. <https://einblogvonvielen.org/tag/nichtbehinderte-menschen/> (Stand: 08.04.2022)
- Framke, Caroline (2021). Netflix's 'Behind Her Eyes' Can't Escape Its Baffling, Deeply Silly Ending: TV Review. <https://variety.com/2021/tv/reviews/behind-her-eyes-review-ending-spoilers-1234909542/> (Stand: 27.04.2023)
- Franzon, Johan (2008). Choices in Song Translation. Singability in Print, Subtitles and Sung Performance. *The Translator* 14(2), 373-399.
- Gambier, Yves (2006). Multimodality and Audiovisual Translation. In: Carroll, Mary; Geryzmisch-Arbogast, Heidrun & Nauert, Sandra (Hg.) *Audiovisual Translation Scenarios: Proceedings of the Marie Curie Euroconferences MuTra: Audiovisual Translation Scenarios – Copenhagen 1-5 May 2006*. https://euroconferences.info/proceedings/2006_Proceedings/2006_Gambier_Yves.pdf (Stand: 02.02.2023)
- Gambier, Yves (2013). The position of audiovisual translation studies. In: Millán, Carmen & Bartrina, Francesca (Hg.) *The Routledge Handbook of Translation Studies*. London: Routledge, 45-59.
- Gambier, Yves & Gottlieb, Henrik (2001). *(Multi)Media Translation*. Amsterdam: John Benjamins.
- Garwood, Ian (2006). The Pop Song in Film. In: Gibbs, John & Pye, Deborah (Hg.). *Close Up 1: Filmmakers' Choices/The Pop Song in Film/Reading Buffy*. London: Wallflower Press, 89-166.
- Gerzymisch-Arbogast, Heidrun & Nauert, Sandra (2005). *MuTra: Challenges of Multidimensional Translation: Proceedings of the Marie Curie Euroconferences*. Saarbrücken: Advanced Translation Research Center.
- Gesetz zur Gleichstellung von Menschen mit Behinderungen (Behindertengleichstellungsgesetz - BGG), BGBl. I 1467, 1468.

- Gesundheit.gv.at (2021). Hörbehinderung & Gehörlosigkeit. <https://www.gesundheit.gv.at/krankheiten/behinderung/taubheit> (Stand: 23.02.2022)
- Greco, Gian Maria (2018). The nature of accessibility studies. *Journal of Audiovisual Translation* 1 (1), 205-232.
- Haig, Raina (2002). Audio Description: art or industry? <http://www.rainahaig.com/pages/AudioDescriptionAorI.html> (Stand: 05.05.2022)
- Halliday, Michael (1978). *Language as Social Semiotic: Social Interpretation of Language and Meaning*. Baltimore: University Park Press.
- Heimlich, Ulrich (2016): *Pädagogik bei Lernschwierigkeiten. Sonderpädagogische Förderung im Förderschwerpunkt Lernen*. Bad Heilbrunn: Julius Klinkhardt.
- Holmes, James (1988 [1972]). The Name and Nature of Translation Studies. *Translated! Papers on Literary Translation and Translation Studies*, 67–80. Amsterdam: Rodopi.
- Holz-Mänttari, Justa (1984). *Translatorisches Handeln: Theorie und Methode*. Helsinki: Suomalainen Tiedakatemia.
- Hurtado, Amparo (1994). Modalidades y tipos de traducción. *Vasos comunicantes* 4, 19–27.
- IMDB (o.J.). Sie Weiß Von Dir. https://www.imdb.com/title/tt9698442/episodes/?ref=tt_ep_epl (Stand: 27.04.2023)
- Incalcaterra McLoughlin, Laura (2018). Audiovisual translation in language teaching and learning. In: Millán, Carmen & Bartrina, Francesca (Hg.) *The Routledge Handbook of Audiovisual Translation*. London: Routledge.
- Jakobson, Roman (1959). On Linguistic Aspects of Translation. In: Brower, Reuben Arthur (Hg.) *On Translation*. Cambridge: Harvard University Press, 232-239.
- Jewitt, Carey, Bezemer, Jeff & O'Halloran, Kay (2016). *Introducing Multimodality*. London: Routledge.
- Jones, Rodney (2013). Multimodal Discourse Analysis. In: Chapelle, Carol (Hg.) *The Encyclopedia of Applied Linguistics*. London: Blackwell.
- Jüngst, Heike (2010). *Audiovisuelles Übersetzen: Ein Lehr- und Arbeitsbuch*. Tübingen: Narr.
- Kaindl, Klaus (2013). *Multimodality and translation*. In: Millán, Carmen & Bartrina, Francesca (Hg.) *The Routledge Handbook of Translation Studies*. London: Routledge, 257-269.
- Kaindl, Klaus (2020). A theoretical framework for a multimodal conception of translation. In: Boria, Monica; Carreres, Ángeles; Noriega-Sánchez, María & Tomalin, Marcus (Hg.) *Translation and multimodality: beyond words*. London: Routledge, 49-70.
- Kress, Gunther & van Leeuwen, Theo (2001). *Multimodal discourse: the modes and media of contemporary communication*. London: Hodder Education.

- Kruger, Jan-Louis (2012). Making Meaning in AVT: Eye tracking and viewer construction of narrative. *Perspectives: Studies in Translatology* 20(1), 67–86.
- Kukkonen, Karin (2014). Plot. In: Hühn, Peter; Meister, Jan Christoph; Pier, John & Schmid, Wolf (Hg.) *Handbook of Narratology*. Berlin: de Gruyter.
- Kurch, Alexander (2018). Produktionsprozesse der Hörgeschädigten-Untertitelungen und Audiodeskription: Potenziale teilautomatisierter Prozessbeschleunigung mittels (Sprach-)Technologien. In: Maaß, Christiane & Rink, Isabel (Hg.) *Handbuch Barrierefreie Kommunikation*. Berlin: Frank & Timme, 437-453.
- Leidmedien (2021). Begriffe über Behinderung von A bis Z. <https://leidmedien.de/begriffe/> (Stand: 30.11.2021)
- Liao, Min-Hsiu (2018). Museums and the creative industries: The contribution of Translation Studies. *Journal of Specialised Translation* 29, 45-62.
- Lorenzo, Lourdes (2010). Subtitling for deaf and hard of hearing children in Spain: A case study. In: Matamala, Anna & Orero, Pilar (Hg.) *Listening to Subtitles: Subtitles for the Deaf and Hard of Hearing*. Bern: Peter Lang, 115-138.
- Maaß, Christiane (2018). Übersetzen in Leichte Sprache. In: Maaß, Christiane & Rink, Isabel (Hg.) *Handbuch Barrierefreie Kommunikation*. Berlin: Frank & Timme, 251-272.
- Maaß, Christiane & Rink, Isabel (2018). Über das Handbuch Barrierefreie Kommunikation. In: Maaß, Christiane & Rink, Isabel (Hg.) *Handbuch Barrierefreie Kommunikation*. Berlin: Frank & Timme, 17-28.
- Mälzer, Nathalie & Wünsche, Maria (2018). Untertitelung für Hörgeschädigte. In: Maaß, Christiane & Rink, Isabel (Hg.) *Handbuch Barrierefreie Kommunikation*. Berlin: Frank & Timme, 327-344.
- Mariotti, Cristina (2015). A Survey on Stakeholders' Perceptions of Subtitles as a Means to Promote Foreign Language Learning. In: Gambier, Yves; Caimi, Annamaria und Mariotti, Cristina (Hg.) *Subtitles and Language Learning*. Bern: Peter Lang, 83–104.
- Maszerowska, Anna (2012). Casting the Light on Cinema: How luminance and contrast patterns create meaning. *MonTI* 4, 65–85
- Mayoral Asensio, Roberto; Kelly, Dorothy & Gallardo, Natividad (1988). Concept of constrained translation: Non-linguistic perspectives of translation. *Meta: Translator's Journal* 33 (3), 356–367.
- Moviepilot (o.J.). Sie Weiß Von Dir. <https://www.moviepilot.de/serie/behind-her-eyes> (Stand: 27.04.2023)

- Musenberg, Oliver (2018). Unterstützte Kommunikation. In: Maaß, Christiane & Rink, Isabel (Hg.) *Handbuch Barrierefreie Kommunikation*. Berlin: Frank & Timme, 361-380.
- Neather, Robert (2012). Intertextuality, Translation and the Semiotics of Museum Presentation: The case of bilingual texts in Chinese museums, *Semiotica* 192, 197–218.
- Netflix (2021a). Barrierefreiheit auf Netflix. <https://help.netflix.com/de/node/116022> (Stand: 30.06.2021)
- Netflix (2021b). Timed Text Style Guide: General Requirements. <https://partnerhelp.netflixstudios.com/hc/en-us/articles/215758617-Timed-Text-Style-Guide-General-Requirements> (Stand: 29.11.2021)
- Netflix (2021c). English Timed Text Style Guide. <https://partnerhelp.netflixstudios.com/hc/en-us/articles/217350977-English-Timed-Text-Style-Guide> (Stand: 29.11.2021)
- Netflix (2021d). French Timed Text Style Guide. https://partnerhelp.netflixstudios.com/hc/en-us/articles/217351577-French-Timed-Text-Style-Guide#h_01EE0KWJ5W0G9XMZJ6Y71SH9SN (Stand: 29.11.2021)
- Netflix (2021e). German Timed Text Style Guide. <https://partnerhelp.netflixstudios.com/hc/en-us/articles/217351587-German-Timed-Text-Style-Guide> (Stand: 29.11.2021)
- Neves, Josélia (2004). Language Awareness through Training in Subtitling. In: P. Orero (Hg.) *Topics in Audiovisual Translation*. Amsterdam: John Benjamins, 127–140.
- Neves; Josélia (2005). *Audiovisual Translation: Subtitling for the Deaf and Hard of Hearing*. Dissertation, University of Roehampton.
- Neves, Josélia (2018). Subtitling for deaf and hard of hearing audiences: Moving forward. In: Pérez-González, Luis (Hg.) *The Routledge Handbook of Audiovisual Translation*. London: Routledge, 82-95.
- Nida, Eugene (1964). *Toward a Science of Translating: With Special Reference to Principles and Procedures Involved in Bible Translating*. Leiden: E. J. Brill.
- O'Hagan, Minako, & Mangiron, Carme (2013). *Game localization: translating for the global digital entertainment industry*. Amsterdam, Philadelphia: John Benjamins
- ÖHTB (2022). Taubblind. <https://oehtb.at/taubblind> (Stand: 23.02.2022)
- ORF.at (2023). Barrierefreiheit im ORF. <https://der.orf.at/kundendienst/service/barrierefrei100.html> (Stand: 14.05.2023)
- O'Sullivan, Carol (2013). Introduction: Multimodality as challenge and resource for translation. *Journal of Audiovisual Translation* 20, 2-14.

- O’Sullivan, Carol & Cornu, Jean-François (2018). History of audiovisual translation. In: Pérez-González, Luis (Hg.) *The Routledge Handbook of Audiovisual Translation*. London: Routledge, 15-30.
- Pérez-González, Luis (2007). Intervention in New Amateur Subtitling Cultures: A multimodal account. *Linguistica Antverpiensa* 6, 67-80.
- Pérez-González, Luis (2015). *Audiovisual Translation: Theories, Methods and Issues*. London: Routledge.
- Pérez-González, Luis (2019). Multimodality. In: Baker, Mona & Saldanha, Gabriela (Hg.) *Routledge Encyclopedia of Translation Studies*. London: Routledge, 346-351.
- Queen, Robin (2004). „Du hast jarkeene Ahnung“: African American English Dubbed into German. *Journal of Sociolinguistics* 8(4), 515-537.
- Reiß, Katharina (1971). *Möglichkeiten und Grenzen der Übersetzungskritik: Kategorien und Kriterien für eine sachgerechte Beurteilung von Übersetzungen*. München: Hueber.
- Reiß, Katharina & Vermeer, Hans (1984). *Grundlegung einer allgemeinen Translationstheorie*. Tübingen: Niemeyer.
- Remael, Aline (2001). Some thoughts on the study of multimodal and multimedia translation. In Yves Gambier & Gottlieb, Henrik (Hg.) *(Multi)Media Translation: Concepts, Practices, and Research*. Amsterdam: John Benjamins, 13–22.
- Remael, Aline (2003). Mainstream Narrative Film Dialogue and Subtitling, *The Translator* 9 (2), 225–247.
- Remael, Aline (2007). Sampling subtitling for the deaf and the hard-of-hearing in Europe. In: Díaz-Cintas, Jorge; Orero, Pilar; Remael, Aline (Hg.) *Media for All: Subtitling for the Deaf, Audio Description, and Sign Language*. Amsterdam: Rodopi, 21-52.
- Remael, Aline & Reviers, Nina (2015). Recreating Multimodal Cohesion in Audio Description: A Case Study of Audio Subtitling in Dutch Multilingual Films. *New Voices in Translation Studies* 13, 50-75.
- Remael, Aline & Reviers, Nina (2018). Multimodality and audiovisual translation. Cohesion in accessible films. In: Pérez-González, Luis (Hg.) *The Routledge Handbook of Audiovisual Translation*. London: Routledge, 260-280.
- Remael, Aline; Reviers, Nina & Vandekerckhove, Reinhild (2018). From Translation Studies and Audiovisual Translation to Media Accessibility: Some Research Trends. In: Gambier, Yves; Pinto, Sara Ramos (Hg.) *Audiovisual Translation: Theoretical and Methodological Challenges*. Amsterdam: John Benjamins, 65-77.

- Richtlinie 2010/13/EU des Europäischen Parlaments und des Rates vom 10. März 2010 zur Koordinierung bestimmter Rechts- und Verwaltungsvorschriften der Mitgliedstaaten über die Bereitstellung audiovisueller Mediendienste, ABl. L 2010/95, 1.
- Rink, Isabel (2018). Kommunikationsbarrieren. In: Maaß, Christiane & Rink, Isabel (Hg.) *Handbuch Barrierefreie Kommunikation*. Berlin: Frank & Timme, 29-66.
- Romero-Fresco, Pablo (2015). *The Reception of Subtitles for the Deaf and Hard of Hearing in Europe*. Bern: Peter Lang.
- Romero-Fresco, Pablo (2018). In support of a wide notion of media accessibility: Access to content and access to creation. *Journal of Audiovisual Translation* 1 (1), 187-204.
- Romero-Fresco, Pablo (2019a). *Accessible Filmmaking: Integrating translation and accessibility into the filmmaking process*. London: Routledge.
- Romero-Fresco, Pablo (2019b). Subtitling for the deaf and hard of hearing. In: Baker, Mona und Saldanha, Gabriela (Hg.) *Routledge Encyclopedia of Translation Studies*. London: Routledge, 549-554.
- Romero-Fresco, Pablo & Dangerfield, Kate (2022). Accessibility as a Conversation. *Journal of Audiovisual Translation* 5 (2), 15-34.
- Royce, Terry (2007). Intersemiotic Complementarity: A Framework for Multimodal Discourse Analysis. In: Royce, Terry & Bowcher, Wendy (Hg.) *New Directions in the Analysis of Multimodal Discourse*. Mahwah, NJ: Lawrence Erlbaum, 63-109.
- Ryan, Marie-Laure (2004). *Narrative across Media: The Languages of Storytelling*. Lincoln: University of Nebraska Press.
- Salzburger Nachrichten (2020). 30 Prozent der Netflix-Inhalte müssen künftig aus Europa kommen: AVMD-Richtlinie beschlossen. <https://www.sn.at/panorama/medien/30-prozent-der-netflix-inhalte-muessen-kuenftig-aus-europa-kommen-avmd-richtlinie-beschlossen-95811589> (Stand: 18.11.2021)
- Schuppener, Saskia & Bock, Bettina (2018). Geistige Behinderung und barrierefreie Kommunikation. In: Maaß, Christiane & Rink, Isabel (Hg.) *Handbuch Barrierefreie Kommunikation*. Berlin: Frank & Timme, 221-247.
- Schwalb, Helmut & Theunissen, Georg (³2018). *Inklusion, Partizipation und Empowerment in der Behindertenarbeit: Best-Practice-Beispiele: Wohnen – Leben – Arbeit – Freizeit*. Stuttgart: Kohlhammer.
- Sell, Anne Catherine (2015). Translation in Japanese Museums: A study of multimodal linguistic landscapes, Dissertation, Melbourne: Monash University.

- Snell-Hornby, Mary (1988). *Translation Studies: An Integrated Approach*. Amsterdam: John Benjamins.
- Snell-Hornby, Mary (2006). *The Turns of Translation Studies: New paradigms or shifting viewpoints?* Amsterdam, Philadelphia: John Benjamins.
- Soffritti, Marcello (2018). Multimodal Corpora and Audiovisual Translation Studies. In Pérez-González, Luis (Hg.) *The Routledge Handbook of Audiovisual Translation*. London: Routledge, 334–349.
- Sternberg, Meir (2009). Epilogue: How (Not) to Advance Toward the Narrative Mind. In: Brône, Geert & Vandaele; Jeroen (Hg.) *Cognitive Poetics: Goals, Gains and Gaps*. Berlin: de Gruyter, 455 -532.
- Stöckl, Hartmut (2004). In Between Modes: Language and Image in Printed Media. In: Ventola, Eija; Charles, Cassily & Kaltenbacher, Martin (Hg.) *Perspectives on Multimodality*. Amsterdam, Philadelphia: John Benjamins, 9-30.
- Szarkowska, Agnieszka (2013). Towards interlingual subtitling for the deaf and the hard of hearing. *Perspectives: Studies in Translatology* 21 (1), 68-81.
- Taylor, Christopher (2004). Multimodal text analysis and subtitling. In: Ventola, Eija; Charles, Cassily & Kaltenbacher, Martin (Hg.) *Perspectives on Multimodality*. Amsterdam, Philadelphia: John Benjamins, 153-172.
- Taylor, Christopher (2016). The Multimodal Approach in Audiovisual Translation. *Target* 28(2), 222–236.
- Tiittula, Liisa & Hirvonen, Maija (2019). Siehst du, was ich höre?: Audiovisuelle Multimodalität aus der Perspektive von Seh- und Hörbehinderten. In: Giessen, Hans; Lenk, Hartmut; Tienken, Susanne & Tiittula, Liisa (Hg.) *Medienkulturen – Multimodalität und Intermedialität*. Bern: Peter Lang, 245-259.
- Titford, Christopher (1982). Sub-titling: Constrained translation. *Lebende Sprachen* 27 (3), 113-116.
- Tseng, Chiao-I (2013). *Cohesion in Film. Tracking Film Elements*. Basingstoke: Palgrave Macmillan.
- Übereinkommen über die Rechte von Menschen mit Behinderungen sowie das Fakultativprotokoll zum Übereinkommen über die Rechte von Menschen mit Behinderungen (UN-BRK), BGBl. III Nr. 155/2008.
- Udo, John-Patrick & Fels, Deborah (2010). The rogue poster-children of universal design: Closed captioning and audio description. *Journal of Engineering Design* 21 (2-3), 207-221.

- USH+TB (o.J.) Taubblindheit. <https://usher-taubblind.at/taubblindheit/> (Stand: 23.02.2022)
- Vandaele, Jeroen (2018). Narratology and audiovisual translation. In: Pérez-González, Luis (Hg.) *The Routledge Handbook of Audiovisual Translation*. London: Routledge, 225-241.
- Vanderplank, Robert (1988). The Value of Teletext Subtitles in Language Learning. *ELT Journal* 42(4), 272–281.
- Van Leeuwen, Theo (2005). *Introducing Social Semiotics*. London: Routledge.
- Vermeer, Hans J. (²1990). *Skopos und Translationsauftrag*. Heidelberg: Abt. Allg. Übersetzungs- u. Dolmetschwiss. d. Inst. für Übersetzen u. Dolmetschen d. Univ.
- Wehrmeyer, Ella (2015). Comprehension of television news signed language interpreters: A South African perspective. *Interpreting* 17 (2), 195-225.
- Weissbrod, Rachel & Kohn, Ayelet (2019). Introduction. In Weissbrod, Rachel & Kohn, Ayelet (Hg.) *Translating the Visual: A Multimodal Perspective*. London: Routledge, 1-12.
- Winter, Linda (2014). *Barrierefreie Kommunikation: Leichte Sprache und Teilhabe für Menschen mit Lernschwierigkeiten*. Hamburg: Diplomica.
- Witzel, Jutta (2018). Ausprägungen und Dolmetschstrategien beim Schriftdolmetschen. In: Maaß, Christiane & Rink, Isabel (Hg.) *Handbuch Barrierefreie Kommunikation*. Berlin: Frank & Timme, 303-326.
- Wünsche, Maria (2022). *Untertitel im Kinderfernsehen: Perspektiven aus Translationswissenschaft und Verständlichkeitsforschung*. Berlin: Frank & Timme.
- Wurm, Svenja (2007). Intralingual and Interlingual Subtitling: A Discussion of the Mode and Medium in Film Translation. *The Sign Language Translator and Interpreter* 1 (1), 115-141.
- Yin Yuen, Cheong (2004). The Construal of Ideational Meaning in Print Advertisements. In: O'Halloran, Kay (Hg.) *Multimodal Discourse Analysis: Systemic Functional Perspectives*. London, New York: Continuum, 163-95.
- Yomma (2021). Audismus. <https://www.yomma.de/glossar/audismus/> (Stand: 30.11.2021)

Abstracts

Die vorliegende Masterarbeit setzt sich mit der Untertitelung für Hörgeschädigte (SDH-Untertitelung) auseinander. Sie untersucht den Einfluss von Situations- und Kontextangaben der deutschen, englischen und französischen Untertitel für Hörgeschädigte auf die filmische Narration am Beispiel der Netflix-Serie „Behind Her Eyes“ (2021) anhand des Analysemodells für audiovisuelle Texte nach Chaume (2004). Translationswissenschaftlicher Ausgangspunkt ist die Theorie der Multimodalität. Die SDH-Untertitelung zählt als Teilbereich der Media Accessibility, weshalb auch grundlegende Aspekte der Barrierefreiheit in der vorliegenden Masterarbeit diskutiert werden. Als letzter theoretischer Punkt wird das Feld der audiovisuellen Translation mit besonderem Augenmerk auf SDH-Untertitelung beleuchtet, bevor Methode und Untersuchungsgegenstand dargelegt werden. Die Analyse zeigte, dass die deutsch- und englischsprachigen Untertitel im Durchschnitt die meisten Situations- und Kontextangaben enthielten. Die französischsprachigen Untertitel beinhalteten eine deutlich niedrigere Anzahl an Informationen, was in weiterer Folge bedeuten könnte, dass hörgeschädigte französischsprachige Rezipient:innen ein höheres Maß an mentaler Anstrengung aufwenden müssen, um die filmische Narration nachzuvollziehen, als deutsch- und englischsprachige. Die filmische Narration könnte also in der deutsch- und englischsprachigen Version einfacher verständlich sein als in der französischsprachigen.

The present master's thesis focuses on subtitles for the Deaf and Hard of Hearing (SDH). It aims to understand how information regarding situation and context contained in SDH impacts the film narration. For this purpose, the intralingual subtitles of the Netflix series "Behind Her Eyes" (2021) are analysed using the analytic model for audiovisual texts by Chaume (2004). The theory of multimodal communication serves as a starting point for the analysis. Furthermore, SDH is a form of media accessibility, which is why basic aspects of accessibility are also discussed in this thesis. As a final theoretical point, the field of audiovisual translation is examined with a focus on SDH, before the method and research subject are presented. The analysis showed that, on average, the German and English language versions of the subtitles contained the most information on situation and context. The amount of information contained in French-language subtitles was significantly lower. Consequently, this could mean that hearing-impaired French-speaking recipients have to expend a higher degree of mental effort than German- and English-speaking recipients to understand the film narration. In conclusion, the film narration could be easier to follow in the German and English language versions than in the French language version.

Curriculum Vitae

Vanessa Kerstin Fellner, BA BA

geboren am 19.09.1996

Ausbildung

Oktober 2018 – heute	Masterstudium Translation
März 2017 – April 2021	Bachelorstudium Sprachwissenschaft
Oktober 2015 – August 2018	Bachelorstudium Transkulturelle Kommunikation

Berufserfahrung

April 2023 – heute	Übersetzerin/Projektmanagerin, Simonfay Translation
Juni 2020 – Februar 2023	Untertitlerin, AUDIO2
April 2017 – Mai 2022	Übersetzerin, Pequerrecho Subtitulación
April 2017 – April 2018	Korrekturleserin, Bezirksblätter Wiener Neustadt

Sprachen- und andere Kenntnisse

Sprachen	Deutsch (Muttersprache) Englisch (verhandlungssicher, C2) Französisch (fließend, C1) Spanisch (Grundkenntnisse)
Software-Skills	MS Office (sehr gute Kenntnisse) FAB Subtitler (sehr gute Kenntnisse) Aegisub (gute Kenntnisse)