

Secară, A. & Chereji, R. (2023). Analysing theatre audio introductions: lexical complexity, gender bias in personal characteristics description and structural elements. *The Journal of Specialised Translation*, 40, 241-267. <https://doi.org/10.26034/cm.jostrans.2023.532>

This article is publish under a *Creative Commons Attribution 4.0 International* (CC BY):
<https://creativecommons.org/licenses/by/4.0>



© Alina Secară, Raluca Chereji, 2023

Analysing theatre audio introductions: lexical complexity, gender bias in personal characteristics description and structural elements

Alina Secară, University of Vienna

Raluca Chereji, University of Vienna

ABSTRACT

Starting from the premise that language choices reflect underlying beliefs and assumptions and offer a mirror reflecting attitudes and potential biases present in wider society, the current article provides a linguistic analysis of 52 English audio introductions produced by five UK professional audio describers for a UK-based theatre. We investigate the extent to which the gender of the described characters plays a role in the linguistic choices describers make when describing visible physical markers such as age and race, or abilities and social status. For this purpose, we propose and use a new corpus-based framework for the analysis of personal characteristics in audio introductions. We enrich these results with an analysis of the language complexity of audio introduction texts in terms of their readability and lexical diversity by applying dedicated linguistic investigation metrics. Our results show varying levels of linguistic complexity consistent with the audience profile and play genre, as well as imbalances in the description of characters' personal characteristics depending on their gender. We also confirm previous findings which highlight a degree of standardisation inherent in audio introductions in terms of structural elements present, and their order of appearance, despite the lack of predefined templates.

KEYWORDS

Audio introductions, audio description, language complexity, gender bias, personal characteristics framework, structural elements.

1. Introduction

As a means to translate the visual into words for blind and visually-impaired audiences, audio description (AD) has emerged as a field of investigation in the last two decades. Numerous accounts have been published regarding AD provision and reception as an access service in specific countries (Orero and Wharton 2007; Di Giovanni 2014), standardisation (Vercauteren 2007), practical handbooks (Taylor and Perego 2022), and technology used for its delivery (Orero and Braun 2010; Szarkowska 2011; Tor-Carroggio 2020). The shift towards integrating access services within performances from the outset, rather than as add-ons, has led to further research into the creation and reception of AD within a universal access framework (Accessible Media Inc. 2014; Fryer 2018; Thompson 2018), as well as within easy language contexts (Arias-Badia and Matamala 2020). AD research within the theatre context focused on the development of guidelines (Fryer 2016) and theoretical frameworks (Mazur 2020), and investigations into its language (Hutchinson *et al.* 2020; Reviers 2018).

In the case of theatre performances, ADs are accompanied by audio introductions, also called introductory notes. While both audio descriptions and audio introductions employ linguistic features and techniques to convey visual information into text (York 2007; Reviers 2014), audio introductions

are narrations provided before the start of a performance, whether live or pre-recorded, and give further descriptions of the physical appearance of characters, the costumes, the setting and the staging to make it easier for blind and visually-impaired audience members to follow and enjoy the audio described performance. While audio introductions can be pre-recorded and made available to the public in advance, this article refers to audio introductions delivered live a few minutes before the start of a performance. The audience therefore cannot control the pace of their delivery, nor can they pause to listen again, so comprehension needs to be immediate. This depends on several factors, as highlighted by existing research in AD, and to a lesser extent, audio introductions: an adequate delivery, which includes pace, intonation and style (Ofcom 2000: 10; Fryer and Freeman 2014; Snyder 2014), a familiar structural organisation which meets audience expectations (Iturregui-Gallardo and Solás 2019; Holsanova 2022), and an appropriate language complexity level (Bernabé and Orero 2020: 65; Mazur 2022).

Although not as rich as that pertaining to audio descriptions, research into audio introductions has started to contribute useful findings to the overall description of access services in general, and audio description in particular, with an overview provided by Romero-Fresco (2022) and guidelines within the ADLAB project (Reviere 2014). Investigations into audio introductions as support for live audio described performances are documented by Di Giovanni (2014) and Romero-Fresco and Fryer (2014) who highlight that these serve to complement, and not replace, ADs. Mazur (2020) and Reviere *et al.* (2021) focus on the functions of audio introductions and identify five overarching functions or text types, the importance and relative weight of which depend on the context of each performance. Iturregui-Gallardo and Solás (2019) investigate the structure of audio introductions in opera; the authors provide an overview of the typical order of text-level structural elements and suggest a potential template.

However, one aspect of audio introductions which has not, to our knowledge, been investigated until now are their linguistic characteristics. This is in contrast to existing research on the linguistic features of ADs (Reviere 2018; Soler Gallego 2018; Perego 2019; Hermosa-Ramírez 2021). It is this gap that the current paper is looking to address first. As audio introductions are delivered orally at the start of a performance and are thus shaped by time constraints, it is essential that they are easily and promptly understood by their intended audience and presented in a familiar and consistent order which meets audience expectations. With this in mind, we analyse the language used in audio introductions in terms of linguistic complexity, particularly lexical diversity and readability levels, as well as the extent to which the language choices for describing characters' physical characteristics display gendered differences. To this, we add results regarding the specific structural elements found in audio introductions, and the extent to which these are applied consistently. Studying these three elements together is relevant as new frameworks, such as the one put

forward by Mazur (2022: 102) for a systematic analysis of ADs, assume availability of data linked to lexical choices, syntax and grammar, as well as cohesion and coherence. Finally, our methodology, and specifically, the resulting framework for investigating personal characteristics in audio introductions, could also be applied to further similar investigations in ADs.

2. Research background

2.1. Linguistic complexity

In addition to standard measures such as mean word length, mean sentence length and part-of-speech (PoS) distribution, various metrics have been devised to measure the linguistic complexity of a text. Each such measure presents its own advantages and limitations, capturing “a slightly different facet of the text analysed”, each potentially contributing “to a better understanding of the characteristics of a text” (McCarthy and Jarvis 2010: 390). The most frequently used lexical diversity (LD) index, measuring how many different words are used in a text, is type-token ratio (TTR) (Templin 1957). TTR computes the ratio of the types, i.e., the total number of different words in a text to its tokens, i.e., the total number of words. The higher the TTR, the higher the LD. However, its main drawback is its sensitivity to text length. TTR decreases when texts are longer, repetitions occur, and the number of new types introduced in a text gradually decreases. In trying to mitigate the influence of text length, certain studies use a standardised TTR calculated on the basis of 1000 words. To avoid the drawbacks of TTR, McCarthy and Jarvis (2007) propose another index, Hypergeometric distribution D (HD-D), which calculates a diversity D score by taking 100 random samples of 35–50 tokens, using the TTR formula and then calculating an average D score. Repeated three times and computing an overall average score, it uses “the probability of drawing (without replacement) a certain number of tokens of a particular type from a sample of a particular size” (2010: 383). To this, McCarthy (2005) adds another measure called the measure of textual lexical diversity (MTLD) which, in addition to not depending on text length in the 100–2000 word range, it also takes text structure into account. MTLD cuts the text into sequences which have the same TTR (set to 0.72) and calculates the mean length of the sequences which have the given TTR.

The complexity of a text is also measured in terms of its readability. Flesch-Kincaid readability tests formulae (Kincaid *et al.* 1975) are employed to generate Flesch Readability Ease scores (from 0 to 100) and Flesch-Kincaid Grade Level scores (reflecting the US education system, from Grade 5 to college graduate). The Flesch Readability Ease score formula uses the average number of words per sentence and the average number of syllables per word to generate a result. Its formula uses no semantic information and cannot determine the complexity or familiarity of a word.

Although research on audio introductions is “disappointingly scarce” (Romero-Fresco 2022: 431), with no in-depth studies regarding their linguistic complexity, a few have been performed for audio descriptions, all using TTR as their main measurement for LD. Perego (2019) offers TTR, lexical density, mean word length and mean sentence length as parameters for the textual analysis of 18 museum ADs. She reports a 51.07% TTR and concludes that ADs “seem more lexically and syntactically complex than expected, with their use of opaque technical terms, heavy adjectival phrases and long sentences” (Perego 2019: 333). Hermosa-Ramírez (2021) applies four formal parameters (mean sentence length, open-class word frequencies, PoS distribution and TTR) to investigate the lexico-grammatical patterns of three Catalan (CAT) opera ADs in the Gran Teatre del Liceu in Barcelona and three Spanish (ES) ADs at the Teatro Real in Madrid. She finds that their “mean sentence length (CAT=21.85; ES=19.32) resembles that of the general language corpora for Catalan and Spanish” and “excessive lexical variation is generally avoided in AD, as the aim is to foster access” (Hermosa-Ramírez 2021: 211), with reported standardised TTR of 39.16% for CAT and 37.10% for ES. Soler Gallego (2018) focuses on English ADs for artworks from four museums and reports a 42.5% mean standardised TTR score and mean sentence length of 17.75. Looking to unveil idiosyncratic linguistic patterns of 39 film Dutch ADs, Reviers (2018) performs a comprehensive lexico-grammatical analysis including PoS frequencies, average length of words, sentences and descriptive units, average AD reading speed and standardised TTR. She concludes that the vocabulary variety in the Dutch AD corpus is relatively low, with a marked degree of word repetitions supported by a reported TTR of 38%. We contribute to these findings with results specifically for audio introductions and, in doing so, add further measurements used for linguistic complexity investigations.

2.2. Language choices and bias

Recent research in translation studies points to a possible presence of bias in the way language is used in various domains, including audiovisual translation. Several studies have investigated descriptions of social categories including race (Bernabo 2021), social class (Ranzato 2018), and disability (Bruti and Zanotti 2018). Gender representations have also been assessed in the context of feminist theory (von Flotow and Josephy-Hernández 2018) and ‘engendering’ approaches (De Marco 2016), intersectionality and stereotyping in US television series (Pérez López de Heredia 2016), and gender imbalances in dubbed programmes for children (Drotner 2018; De Ridder 2020), among others.

Our own investigation of potential gender bias in descriptions of personal characteristics in audio introductions is informed by the recent VocalEyes Describing Diversity Report (Hutchinson *et al.* 2020). Following a corpus analysis of 26 audio introductions from VocalEyes’ archive across multiple genres, the authors identify several markers of unconscious bias, as well as

“imbalances and avoidance of” descriptions of personal characteristics across four main categories: race, gender, disability, and age (Hutchinson *et al.* 2020: 13–4). Descriptions of race and gender were particularly reflective of this dynamic, with greater lexical richness displayed in descriptions of white versus black characters, and female versus male characters, which the authors suggest represents “a legacy of language and literary tradition” (Hutchinson *et al.* 2020: 14). Not included as a separate category in their corpus analysis were value judgements, although the authors acknowledge that descriptions across the four above-mentioned categories may reflect or reinforce judgement; they also advise against making value judgements in their twelve principles for describing human characteristics in audio introductions, unless this approach is a deliberate strategy by the describer, which should then be explicitly stated (Hutchinson *et al.* 2020: 61).

Current professional guidelines on audio introductions/descriptions also recognise the importance of race, gender, disability, and age for character descriptions. In the UK, the Ofcom ‘Best practice’ guidelines identify dress, physical characteristics, facial expressions, body language, ethnicity, and age as potentially significant when describing characters (Ofcom 2021: 6), but recommend avoiding subjective language unless relevant to the plot (Ofcom 2021: 6–7). Netflix guidelines require that describers prioritise characters’ “most significant identity traits” in descriptions of visual attributes (Netflix 2020). They recommend that characters’ physical characteristics, which include race, age, or disability markers, be described as factually as possible, with describers avoiding assuming characters’ racial, ethnic, or gender identity (Netflix 2020). This corresponds with current research into describer subjectivity (Bardini 2017; Soler Gallego and Luque 2018; Soler Gallego 2019: 232; Chmiel and Mazur 2021: 560). Although it is generally acknowledged that audio descriptions cannot, and perhaps should not, achieve complete objectivity, a ‘doomed’ pursuit according to Fryer (2016: 172), existing guidelines recommend limiting describers’ appraisals to avoid restricting users’ interpretative freedom (Perego 2019: 343). We enrich this line of research by offering a framework for assessing how personal characteristics are described, and we put forth an applied investigation into the presence of bias in a corpus of audio introductions.

2.3. Structural organisation in audio introductions

Existing research on audio introductions indicates there is little to no standardisation in their structural or thematic organisation, with differences in the type of information included and its order of appearance. This prescriptive structural freedom makes audio introductions a valuable ground for experimentation and investigation of the artistic value of accessible services (Romero-Fresco 2022). In recognising this variation and the lack of a standardised template, Romero-Fresco and Fryer liken audio introductions to jigsaws, which describers can adapt as needed (2014: 17).

In their case, the main elements to include are visual details on main characters and locations, cinematography, storytelling, and plot, with their structural organisation determined by the genre of the performance, as well as the corresponding AD and how this has been developed (Romero-Fresco and Fryer 2014: 18). Di Giovanni (2014) also notes this structural variation in her study on the reception of film audio introductions and descriptions in Italian, but identifies a distinct pattern in how these are organised. This structural pattern consists of an overall film presentation, its genre and structure, its synopsis, visual style descriptions, characters, locations, and cast and production details (Di Giovanni 2014: 2).

Similar structural categories are reported in a study by Reviers *et al.* (2021). Drawing on their earlier corpus-based qualitative analysis of 52 multilingual audio introductions (Reviers and Remael 2013; Reviers and Roofthoof 2018), they identify several recurring categories including general information about the event and performance, production and credits, plot information, scenography, and characters (Reviers *et al.* 2021: 75). Similarly to Romero-Fresco and Fryer (2014), they find that the order of appearance of these structural elements is determined by the performance at hand, though they note these variations appear not only across geographical regions, but also within the work of single describers. Iturregui-Gallardo and Solás (2019) attempt to limit these variations to ensure greater clarity and user-friendliness for users of audio introductions. Drawing on Romero-Fresco and Fryer (2014), as well as prior research by Reviers (2014) and Remael *et al.* (2014), they propose a standardised template for the organisation of opera audio introductions, which consists of similar categories as above, with the addition of an uncategorised element where additional aspects such as textual references can be explained (Iturregui-Gallardo and Solás 2019: 227). We supplement the above studies by offering both quantitative data revealing specific structural elements found in audio introductions, and an analysis of their chronology.

3. Study design

3.1. Research questions

We investigate the linguistic characteristics of audio introductions in terms of their complexity, as well as the presence of physical characteristics descriptions by gender. We also assess the structural features of audio introductions and their order of appearance. As such, the research questions we set out to investigate are:

RQ1: What is the linguistic complexity of audio introductions?

RQ2: Are descriptions of personal characteristics in audio introductions gender biased?

RQ3: Is there consistency in the type and order of appearance of structural elements in audio introductions?

3.2. Corpus

Our corpus consists of 52 live-performed audio introductions written by five female freelance audio describers for a UK-based theatre and include productions staged between 2002 and 2020. The choice of audio describers reflects the real-life working practices of the particular theatre surveyed, which is why male or non-binary describers were not included in the present study. The texts were provided by the describers as Microsoft (MS) Word and non-editable PDF files. We used the Adobe Acrobat PDF to Word tool to convert the audio introductions into editable MS Word files, then converted all texts into .txt files with UTF-8 encoding for lexical and character-based analysis.

The texts in our corpus pertained to performances aimed at the average adult theatregoer and included a range of genres, such as comedies, dramas and variety entertainment. Our corpus also included seven audio introductions (13.46%) pertaining to productions for children; we consider these as a separate sub-corpus to allow for a comparative assessment of potential differences against the general corpus. Table 1 shows the total word count across the two corpora, as well as the mean and standard deviation values serving as the starting point for RQ1.

	Number	Word count	Word M	Word SD
Total corpus	n=52	43,154	825.06	324
Sub-corpus	n=7	3403	584.57	289.1

Table 1. Corpus word-based statistics.

4. Methodology

4.1. Linguistic complexity measures

Measuring how complex and comprehensible a text is requires the use of different parameters. We apply frequently used measures in analysing lexical and textual features to describe the style of the audio introductions and infer their complexity levels as follows: mean word length, mean sentence length, Flesch-Kincaid readability and LD via TTR, HD-D and MTLT.

Mean word length, calculated in characters, provides information on the nature of words used and their perceived difficulty. Mean sentence length provides information on the average number of words per sentence, a

measure also used in readability measurements. Perego cites studies indicating that “for English readability tables show that 8 words or less are considered very easy to read, 11 words easy, 14 words fairly easy, 17 words standard, 21 words fairly difficult, 25 words difficult and 29 words or more very difficult” (2019: 339). In a newspaper English corpus analysis Hearle (2007) reports mean word length of 5.1 characters. In a study investigating the impact of sentence length on the readability of web content when using screen readers, a technology frequently used by the blind, Kadayat and Eika find “a significant effect of sentence length and most participants exhibit the highest comprehension and lowest workload with sentences comprising 16–20 words” (2020: 261).

Following a comparison of LD metrics, McCarthy and Jarvis (2010) recommend using a mix of indexes such as MTLD and vocd-D (or HD-D) rather than any single index, as each approach may help address the question under investigation. Given the above analysis, we include TTR, MTLD and HD-D as measures for LD. In calculating the above, we use the GitHub Python Textstat (Bansal 2016) and the GitHub Python Lexical Richness (Shen 2018) libraries. To assess whether linguistic complexity varies across genres, we compute scores for these metrics for the entire corpus and compare them to scores for the sub-corpus of audio introductions for productions for children.

4.2. Personal characteristics framework

We developed a custom corpus-based framework for assessing personal characteristics in audio introductions based on physical characteristics analyses in the Describing Diversity Report (Hutchinson *et al.* 2020: 13–4), and prior research on subjective description in AD by Mazur and Chmiel (2012). Our framework (Figure 1) consists of three major groups based on gender: ‘Male’, ‘Female’, and ‘Other’ — the latter referring to characters with no gender signifiers, or to non-human characters. While theatre professionals consulted as part of the wider Describing Diversity project do recognise the importance of non-binary gender representations and gender-neutral descriptions (Hutchinson *et al.* 2020: 42), gender-related corpus findings in the report itself focus on representations of maleness and femaleness (Hutchinson *et al.* 2020: 13). This in turn informed our choice for the three gender-based groups surveyed as part of our personal characteristics framework, as did the fact that there were no other gender identities described in our corpus. Each group includes seven parent categories, eleven child categories and six child sub-categories against which we label words or strings of words used in our corpus to describe personal characteristics — we call the labelled results ‘references’.

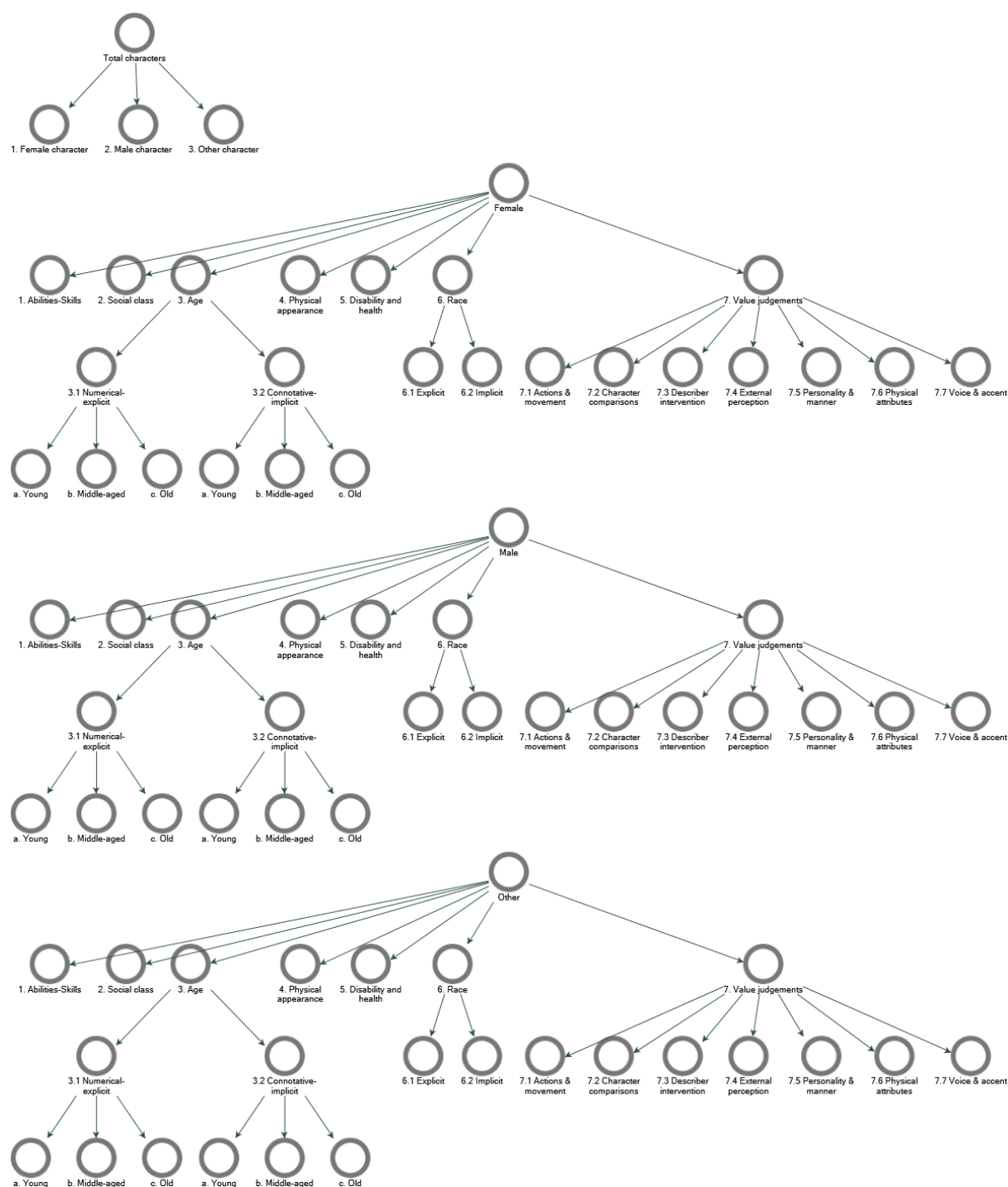


Figure 1. Framework for labelling personal characteristics in audio introductions.

In applying the above framework to our corpus, we define 'Abilities-Skills' in terms of characters' activities, relating both to work and leisure, as well as their aptitudes and interests. We consider specific actions ("playing solitaire", *Enjoy*), and more general depictions of characters' roles ("the women work", *Barnbow Canaries*). The 'Social class' parent category also includes work-related descriptions, but consists specifically of named professions ("the school cleaner", *Into the Woods*; "virologist", *Villette*) and

references to characters' social ranking and status ("the Duke of Cornwall", *King Lear*; "a society lady", *Anna Karenina*).

The 'Age' parent category we divide into two child categories: 'Numerical-explicit' and 'Connotative-implicit'. The former includes explicitly stated descriptions of age, whether through an exact number ("eleven years old", *Annie*), a numerical range ("in his thirties", *The Girl on the Train*) or age groups ("middle aged women", *Romeo and Juliet*). The 'Connotative-implicit' child category refers to age as inferred through descriptions of age-related physical transformations ("white hair", *Hamlet*; "receding grey hair", *Doctor Faustus*). Both child categories for age are subdivided into three child sub-categories corresponding to the young, middle-aged and old age groups.

For 'Physical appearance', we consider traditional descriptions of characters' outward image and include references to dress and clothing ("in sweatshirts, tracksuit bottoms", *Beryl*), hair ("long fair hair", *Chitty Chitty Bang Bang*), and physique ("rather thin", *Hamlet*). As per the Describing Diversity Report (Hutchinson *et al.* 2020), we include 'Disability and health' and 'Race' as further parent categories, with the latter divided into two child categories: 'Explicit', or stated descriptions of race and ethnicity ("black", *Little Voice*; "Nigerian", *Barbershop Chronicles*), and 'Implicit', or implied race based on racial markers ("an afro hair style", *Romeo and Juliet*; "almost white, skin", *Great Expectations*).

For 'Value judgements', we include markers of subjective evaluations of the characters and their attributes. These consist primarily of adjectival and adverbial constructions serving to interpret *how* a character acts ("forthright manner", *King Lear*) or appears ("solemn, but pretty, face", *Barnbow Canaries*). We subsequently identify seven child categories, broad themes for value judgement references, which we discuss in section 5 (*Results and discussion*).

Our 52 audio introductions were imported into NVivo, a qualitatively data analysis software (QSR 2020) as individual .txt files. Each file was parsed individually and each reference identified was manually labelled to the relevant gender and parent or child (sub-)category. References were restricted to the smallest textual unit conveying a single idea; in some cases, this entailed coding individual words ("tall", *Romeo and Juliet*), while in others, it required including entire phrases as they related to the same feature ("plays a guitar with some humour", *Uncle Vanya*). Enumerations were divided into separate references where they described disparate features of a character ("tall / slim", *Hamlet*), but grouped when presenting a unitary description ("a wide, white headband and a loosely fitting white dress with long sleeves", *A Christmas Carol*). Each reference was added as many times as it appeared in a given file, and cross-categorisation was allowed where references pertained to more than a single category (e.g.,

"a bit of a mad cap inventor", *Chitty Chitty Bang Bang*, in both 'Abilities and Skills' and 'Value judgements').

Once the coding was finalised, all references were analysed per gender, parent and child (sub-)categories supported by the metadata and automatic counts in NVivo.

4.3. Structural organisation framework

We drew on the structural categories of audio introductions identified by Di Giovanni (2014), Iturregui-Gallardo and Solás (2019) and Reviers *et al.* (2021), and developed our own structural organisation framework comprised of eight major and eleven subordinate structural elements: 'General event information' ('Descriptors', 'Venue', 'Length', 'Announcements'), 'General presentation of the piece' ('Genre', 'Structure', 'History', 'Creative team'), 'Synopsis', i.e., plot information, 'Stage directions', 'Visual style', i.e., scenography, 'Characters' ('Overall cast', referring to ensembles or undefined character groups, 'Individual characters', 'Actors'), 'Locations', and 'Production information', i.e., credits. Since three major structural elements ('General event information', 'General presentation of the piece' and 'Characters') included subordinates, we considered these individually in our analysis, as opposed to only the major element. Therefore, in total, our framework contains 16 structural elements (Figure 2).

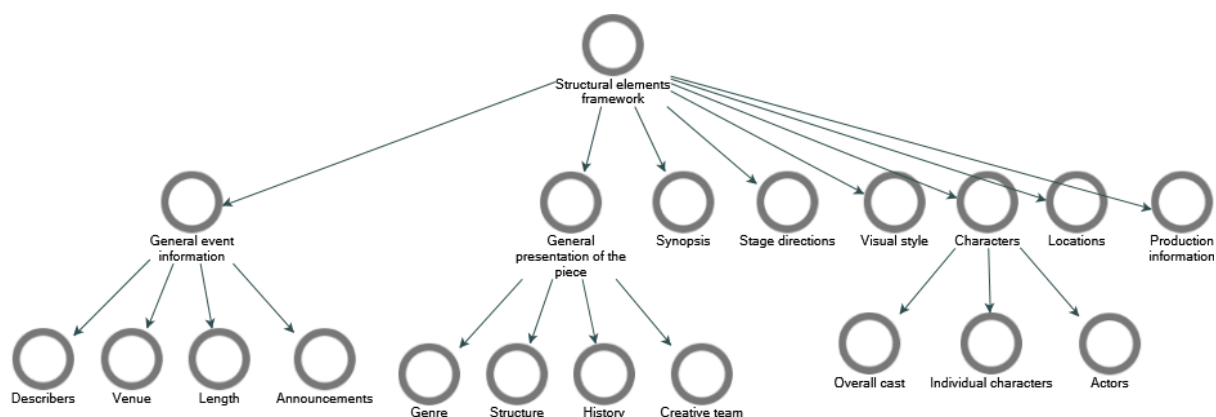


Figure 2. Framework for mapping structural elements in audio introductions.

We parsed each audio introduction individually against our framework and recorded the presence and order of appearance of each structural element. Where one structural element appeared multiple times across a single text, we included only its first appearance in our analysis. We then counted the number of times each structural element appeared across all files and calculated its frequency of appearance as a percentage relative to the entire corpus. Using the recorded numerical position of each element, we established its order of appearance by calculating its mean across the 52 audio introductions in the corpus.

5. Results and discussion

Following our corpus analysis, we present our findings and discussion below, in response to our three research questions.

5.1. Linguistic complexity

	Word length	Sentence length	TTR	HD-D	MTLD	FK Ease	FK Grade
Total corpus (n=52)	M 4.55 (SD 0.18)	M 19.18 (SD 4.15)	M 0.50 (SD 0.05)	M 0.85 (SD 0.02)	M 90.49 (SD 21.35)	M 76.2 (SD 8.1)	M 7.6 (SD 2.1)
Sub-corpus (n=7)	M 4.31 (SD 0.16)	M 18.46 (SD 5.16)	M 0.49 (SD 0.06)	M 0.82 (SD 0.02)	M 74.24 (SD 10.6)	M 83.83 (SD 7.6)	M 6.2 (SD 2.2)

Table 2. Lexical complexity measurements.

For the total corpus, we record higher LD values for TTR 0.50 (SD=0.05) compared to previous studies (Reviere 2018; Soler Gallego 2018; Hermosa-Ramírez 2021), pertaining to an average LD. Our mean word length is slightly higher (4.55, SD=0.18) than that recorded in Perego (2019) (4.39, SD=2.26), but remains in the average complexity category. The mean sentence length in Perego (2019) was recorded at 19.32 (SD=7.87), a very similar value to ours 19.18 (SD=4.15), indicating that slightly complex morpho-syntactic structures are used in audio introductions. Hermosa-Ramírez (2021) reports 21.85 for Catalan and 19.32 for Spanish and Soler Gallego (2018) 19 for English.

When measuring our total corpus against our sub-corpus of audio introductions for children's performances for genre-based differences, the LD scores for two measures, TTR and HD-D, suggest only a negligible difference between the two. However, the MTLD, which was shown not to correlate with text length, shows a clear difference in LD between the two corpora, with the sub-corpus displaying a lower mean 74.24 (SD=10.60) compared to 90.49 (SD=21.35). This genre difference is also supported by the FK Grade 7.6 (SD=2.1), i.e., fairly easy to read, for the total corpus, and 6.2 (SD=2.2), i.e., easy to read, for the sub-corpus. The mean word lengths values are very similar at 4.55 (SD=0.18) and 4.31 (SD=0.16), both indicating average length. At 18.46 (SD=5.16), the sub-corpus mean sentence length values point to average structures, while the total corpus values (19.18, SD=4.15) indicate more complex morpho-syntactic structures. However, according to Kadayat and Eika (2020), even at the upper limit, this range should still be manageable for the average user.

In response to our RQ1, the differences in the values recorded for our two corpora regarding mean sentence length, FK scores and MTLD indicate that audio describers are adjusting the language they employ, depending on genre and audience profile. Moreover, the FK score for the total corpus, the LD scores, and the mean word length values point to a fairly easy to read level as defined by the FK parameters and support the idea that the language used in our audio introductions is carefully considered by the describers to grant comprehension to the general public. Similar to Hermosa-Ramírez (2021), we acknowledge a certain difficulty in comparing our results with those reported in other studies, given differences in methodology, languages or indexes used; nonetheless, we notice certain common traits regarding language usage. As TTR values rarely go above 50% and mean word length and readability scores do not show complexity, we observe that audio describers adapt their language choices to match an average language competence level suitable for an access context. However, one area where there seems to be a tendency towards more complexity is syntax; as in Perego's study (2009), we also record a higher than average mean sentence length.

5.2. Personal characteristics

We consider male (M) and female (F) characters in the following analysis, as per the Describing Diversity Report (Hutchinson *et al.* 2020). Figure 3 presents the total personal characteristics references coded by gender across the seven parent categories in our framework.

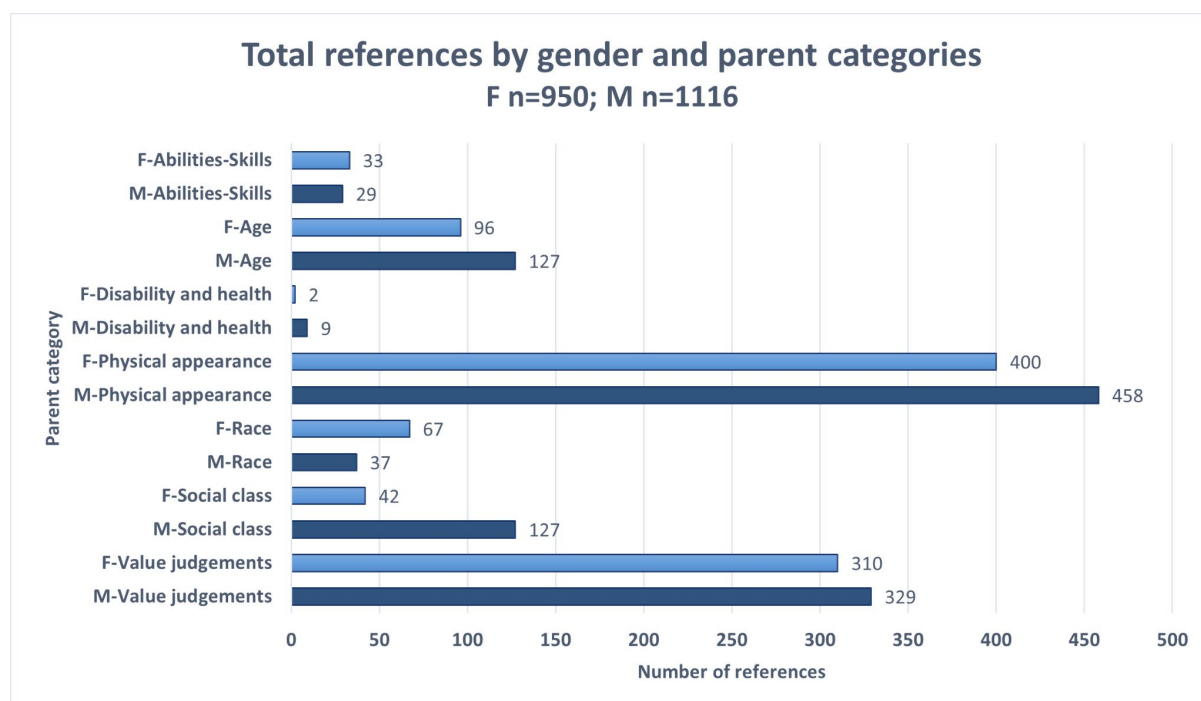


Figure 3. Total corpus references by gender and parent categories.

We identify 270 male and 186 female characters in our corpus. Out of the total 2066 references collected, 1116 pertain to males and 950 to females. This suggests that although there are 18.43% more male than female characters in our corpus, the number of labelled references shows only an 8.03% difference in the number of male references to female references.

There are, however, imbalances in the nature and type of references per gender. Regarding 'Age' references, we note an uneven distribution of age groups by gender. In the dominant 'Numerical-explicit' child category, we find the same number of young and middle-aged male characters ($n=45$; 1:1 ratio), but twice more young than middle-aged female characters ($n=52$ and $n=26$; 2:1 ratio), an asymmetry likely attributable to character distribution in each play. Gender representation is more balanced for 'Connotative-implicit' descriptions, except for the 'Middle-aged' child sub-category, where male characters have 70.37% more references than female characters. These descriptions centre on male-specific depictions of ageing ("receding grey hair", *Dr Faustus*; "silver grey hair and beard", *Around the World in 80 Days*), producing more numerous and detailed descriptions than for female characters, which focus on hair alone ("grey hair", *Enjoy*; "grey streaked hair", *Strictly Ballroom*). Figure 4 outlines age references by gender across all child (sub-)categories.

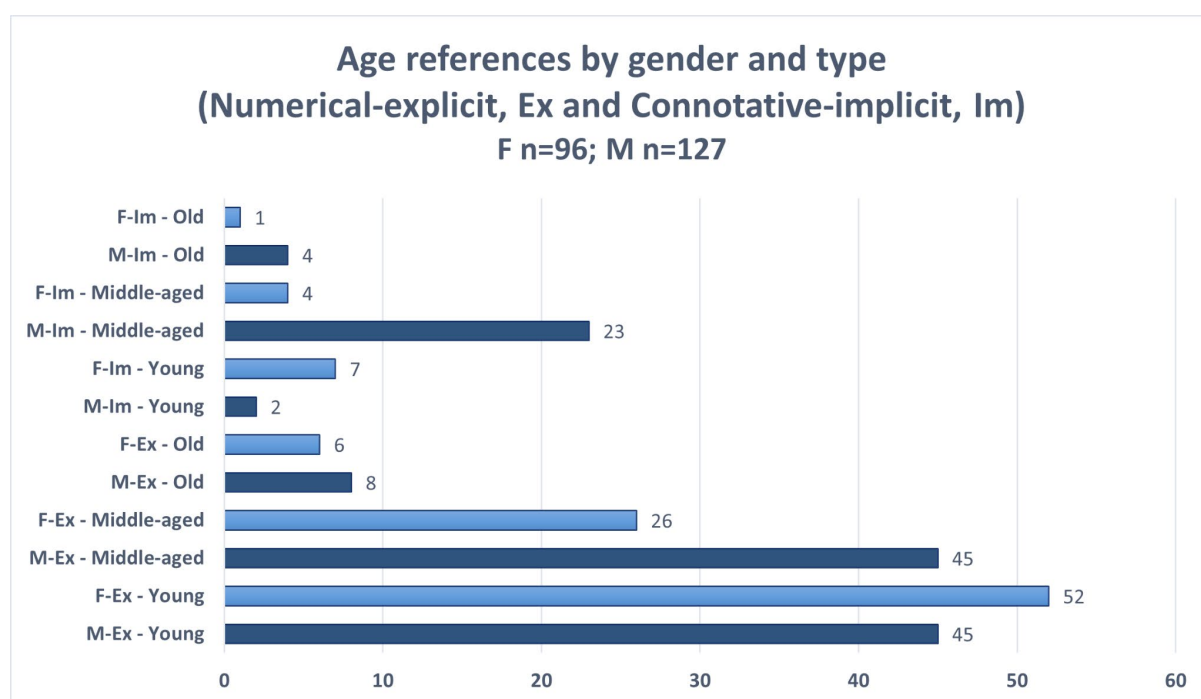


Figure 4. Age references by gender and type (Numerical-explicit and Connotative-implicit).

We identify similar gender-based discrepancies for the 'Abilities-Skills' and 'Social class' parent categories and present these jointly given the link between characters' skills and abilities, and their social standing. While 'Abilities-Skills' references are similar across genders (M $n=29$; F $n=33$), we identify three times more references to male versus female social class

(M n=127; F n=42). This translates into a much greater richness and variety of male social roles: royalty ("King of Scotland", *Macbeth*), military ("the officers", *Our Country's Good*), clergy ("clergyman", *Pride and Prejudice Sort Of*), politics ("Chief Whip", *This House*), trades ("London lawyer", *Great Expectations*) and artists ("rock star", *Doctor Faustus*). This variety also concerns abilities, producing well-rounded, nuanced male characters with diverse interests and creativity. In contrast, female character descriptions are more limited: many 'Abilities-Skills' references pertain to work, which is either vague or nondescript ("works in the Proctor's house", *The Crucible*) or related to housework ("bustles around doing housework", *Enjoy*). The few positions of power included are also general ("the head of the lab", *Villette*; "in charge of the Orphanage", *Annie*). This correlates with descriptions of female social roles, which include few professions beyond the domestic ("the housekeeper", *Pygmalion*), junior ("assistant desk clerk", *The Graduate*), or traditionally female roles ("Juliet's nurse", *Romeo and Juliet*). For these two categories however, we suggest that these discrepancies stem from the source material itself, rather than describers' choices. Many of the texts in our corpus include very few female characters (e.g., *Macbeth*), and those present often have a limited social status, in keeping with the particular historical reality depicted in the original texts.

The 'Physical appearance' parent category reveals a disproportionate representation of female over male physical characteristics. We identify only 6.75% more descriptions of male over female physical appearance, which does not reflect the 18.43% difference between the number of male and female characters present overall in our corpus. This suggests that although there are fewer female characters in our corpus, their physical appearance, in terms of number of references collected, is almost as amply described as males'. Descriptions broadly cover three main themes for both genders: hair and hairstyles ("long straight brown hair", *Hamlet*), face and body ("medium build and height"), and clothing ("wears a pale dress", *A Christmas Carol*). However, female hair descriptions are more elaborate and detailed than for male characters ("her red hair is lacquered into a neat curly bob", *The Graduate*; "streams of loose hair tumbling down her back", *Random*; "light ginger hair which she wears in a loose bun, tendrils of hair falling around her face", *Uncle Vanya*). These results are in keeping with findings in the Describing Diversity Report, where women's bodies were described "in more richness and more at length than those of men", likely reflecting literary tradition and language legacies on one hand (Hutchinson *et al.* 2020: 14), and imbalanced representations on the other. Our analysis also shows instances of overlap with 'Value judgements' references for female characters, with descriptions occasionally accompanied by evaluative or euphemistic qualifiers ("a leather belt emphasizing her large behind", *Caucasian Chalk Circle*; "wears a full skirt with a saucy red underskirt peeping out from beneath", *Sherlock*).

Related to physical appearance are depictions of 'Disability and health', and 'Race'. The former includes 63.63% more references for male than female characters. Although the total number of references across genders is itself quite small (N=11), we note a slight nuance for female descriptions ("bowed over her walking stick", *The Crucible*) compared to those for males ("He uses a walking stick", *Enjoy*), which are generally more factual. We also find gendered differences for 'Race' (Figure 5). Overall, there are 28.85% more references to female racial characteristics compared to males, with a 41.87% difference for implicit references. The descriptions, however, are similar and suggest race by describing hair ("an afro hair style", *Romeo and Juliet*; "reddish hair", *Anna Karenina*), skin tone ("dark-skinned", *Barnbow Canaries*; "pale of face", *Into the Woods*) and sometimes accent ("English accent", *Sunshine on Leith*). We also note a preponderance of distinctly Caucasian features such as blonde or red hair and fair complexions, particularly for female characters, compared to more vague descriptions of male race ("dark", *Villette*; "light brown hair", *Mary Shelley*). We identify twice more explicit race references (2:1) for male characters than for females, though in both cases these are expressed through explicit mentions of race or ethnicity ("Yugoslavian", *Loserville*; "Negro", *The Crucible*).

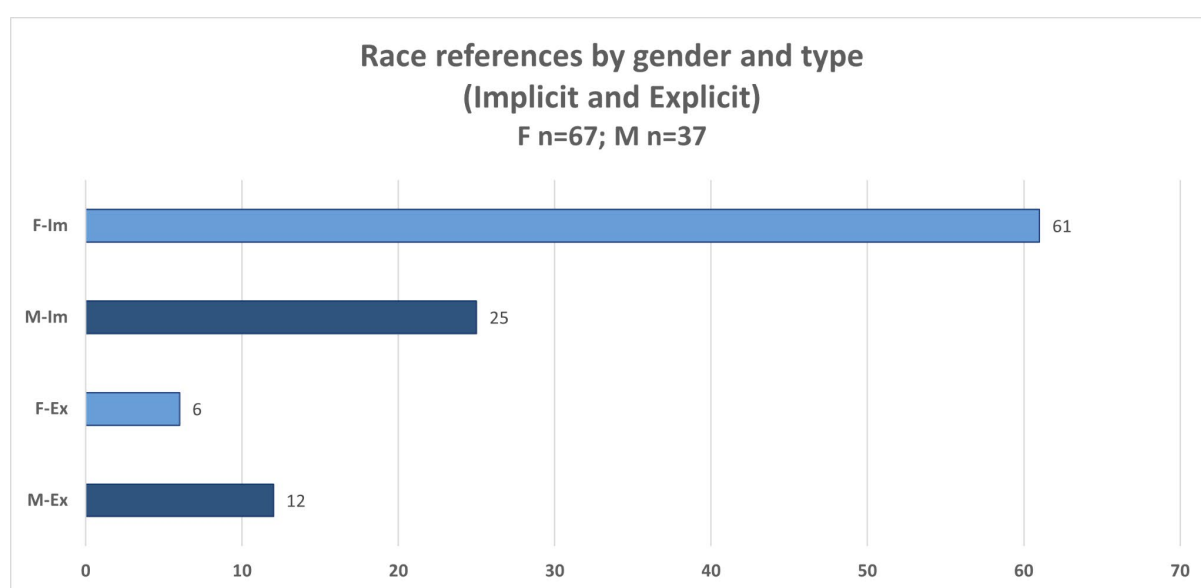


Figure 5. Race references by gender and type (implicit and explicit).

We note the same trend for 'Value judgements' as for 'Physical appearance'. We coded only 2.97% more references for male versus female characters, despite having 18.43% more male than female characters in the total corpus. This suggests that while evaluative language is present across genders, it is not evenly distributed, with proportionally more value judgments for female characters. Regarding the descriptions themselves, these broadly fit within seven themes for both genders (Figure 6). Representation across these themes is mostly balanced, save for 'Personality & manner' and 'Physical attributes', where we note cross-gender differences. There are 20.73% more references to male

characters' 'Personality & manner', but 14.95% more evaluative descriptions of female 'Physical attributes', despite fewer female characters in the corpus.

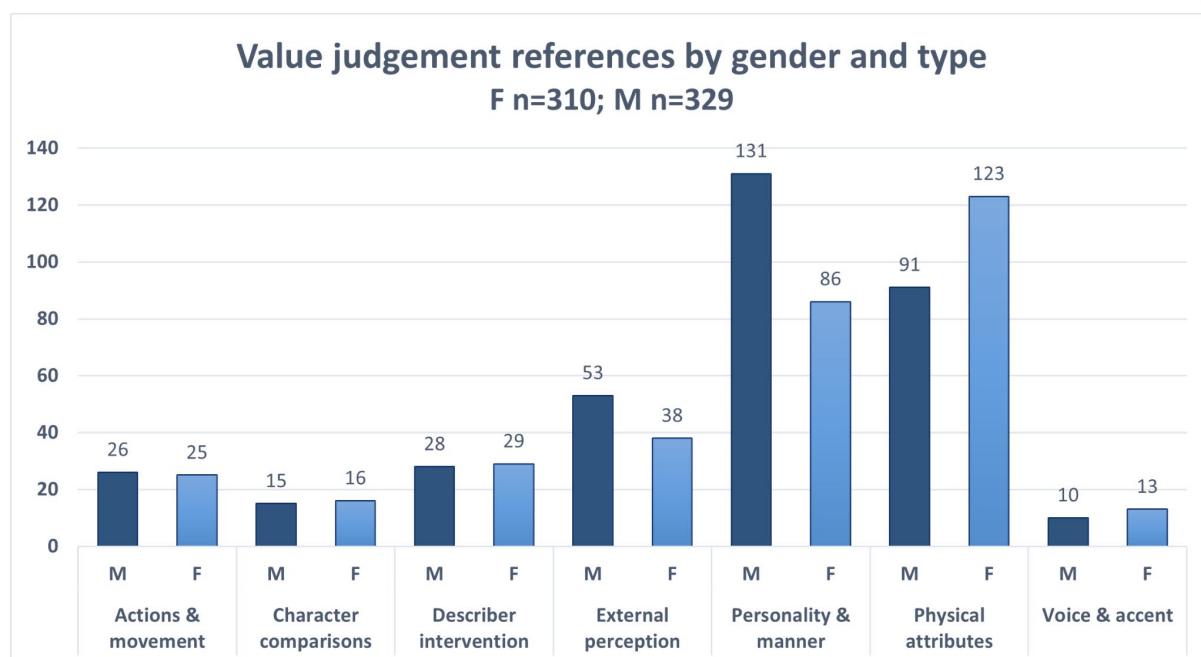


Figure 6. Value judgement references by gender and type.

Another notable difference concerns the 'Descriptor intervention' theme. References to male characters tend to include hedging devices ("He might be described as geeky in his pursuit of knowledge", *Doctor Faustus*; "Somewhat unconvincing in the role of chauffeur", *Enjoy*), which serve to qualify the weight of the descriptions and create distance from the characters. In contrast, interventions for female characters tend to be more emphatic and categorical. Here descriptions directly address the characters' physical appearance ("generous hips which she uses to great effect", *Blues in the Night*; "While not fat she no longer has the sleek body of her youth", *Little Voice*) and call into question their drives and motivations ("reflecting Alice's state of mind/ lucidity?", *Still Alice*; "Is this confidence just skin deep?", *The Graduate*). In other instances, female characters are described using established imagery and stereotypes: whether in popular, folkloristic terms ("is in the tradition of Fairy Stories, old and ugly", *Into the Woods*), or familiar tropes ("agony aunt", *Pride and Prejudice Sort Of*; "Nora Batty style!", *Uncle Vanya*).

In response to RQ2, our analysis reveals notable differences between descriptions of male and female characters. In some instances, such as the 'Numerical-explicit' child category and the 'Abilities-Skills' and 'Social class' parent categories, these imbalances stem from the fact that describers were probably constrained in terms of *what* to describe by casting choices and the source material itself, which reflected social and gender dynamics pertaining to the historical reality in question. In other cases, however, imbalances may result from *how* personal characteristics are described.

These gendered differences are particularly salient in implicit depictions of age and race, and in physical appearance descriptions, which tended to be longer and richer in detail for female characters compared to those of males, consistent with findings in the Describing Diversity Report (Hutchinson *et al.* 2020: 14). We also note describers tend to “adopt the tone of the original production” and “slip into judgements” (Hutchinson *et al.* 2020: 61), particularly in their interventions concerning female characters.

5.3. Structural elements and their order of appearance

Assessed against our framework, we found that the most prevalent structural element in our corpus, whether major or subordinate, was visual style, with all 52 audio introductions containing descriptions of visual setting, such as stage design and lighting. The same results were recorded for the ‘Individual characters’ sub-category, with 100% of the texts surveyed containing references to specific characters in the performance.

The second most common subordinate structural element was ‘Announcements’ (90.38%), followed by ‘Venue’ (84.62%); both pertain to ‘General event information’, the second most frequent major structural element overall, after ‘Visual style’. Also present across most texts were references to the creative ensemble (‘Actors’, 84.62%; ‘Creative team’, 82.69%), although only half of the audio introductions surveyed (50%) included ‘Production information’, i.e., details about the production credits. Common were also mentions of the ‘Describers’ themselves (78.85%), as well as information about ‘Synopsis’, i.e., the plot of the performance (73.08%).

On the other hand, references to ‘Genre’ of the play and ‘Locations’, i.e., where the action took place, were both only present in 26.92% of audio introductions; they represent the least common subordinate, and respectively major, structural element in the framework. Also infrequent were references to ‘Stage directions’ (32.69%), the second least common major structural element. Figure 7 outlines the frequency with which each structural element in our framework was present in our corpus.

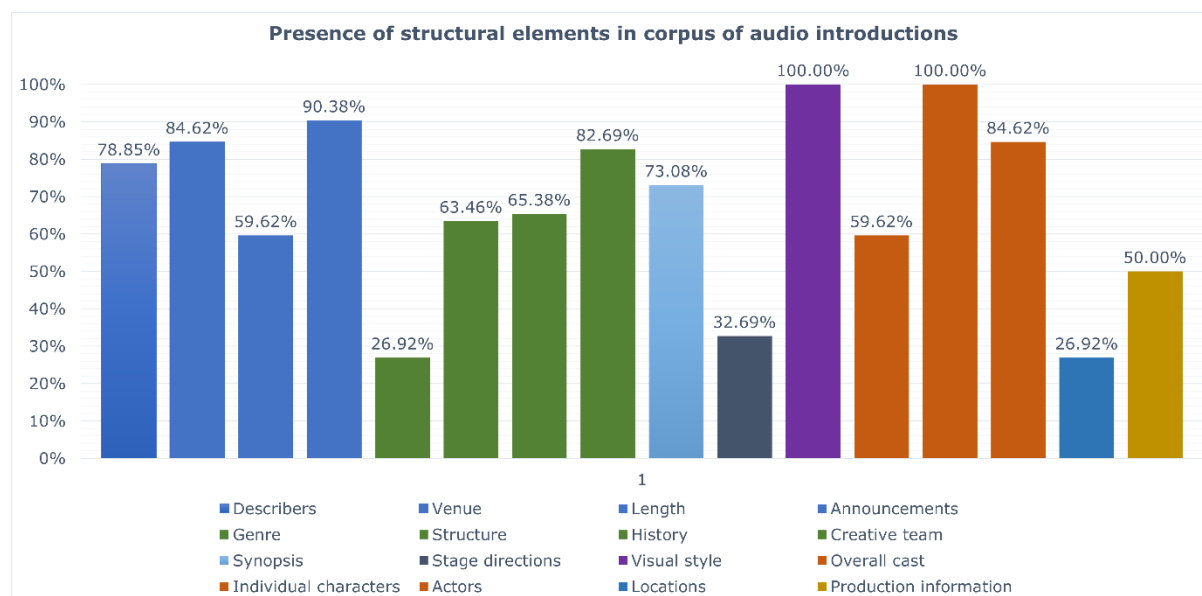


Figure 7. Presence of structural elements in our corpus of audio introductions.

Regarding order of appearance, we found that most audio introductions began on average with references to 'Venue' ($M=1$), followed by 'Describers' ($M=2.097$). Descriptions of 'Genre' ($M=3.071$), 'Production information' ($M=3.884$), 'History' ($M=4.352$), 'Creative team' ($M=5.139$) and 'Synopsis' ($M=5.342$) were presented, on average, in succession during the first half of the introductions, and occupied positions 3 to 7 in terms of order of appearance. Positions 8 to 13 included both practical information about the performance and elements of scenography and characters, with small differences in order of appearance across six structural elements ('Structure', $M=7.09$; 'Length', $M=7.354$; 'Visual style', $M=7.423$; 'Individual characters'; $M=7.442$; 'Overall cast'; $M=7.451$; 'Locations', $M=7.5$). Antepenultimate position 14 tended to introduce the 'Actors' ($M=8.636$) in the performance. Announcements' ($M=9.276$) and 'Stage directions' ($M=10.352$) tended to be in the penultimate and respectively, final position. Figure 8 illustrates the average order of appearance for all 16 structural elements across our corpus.

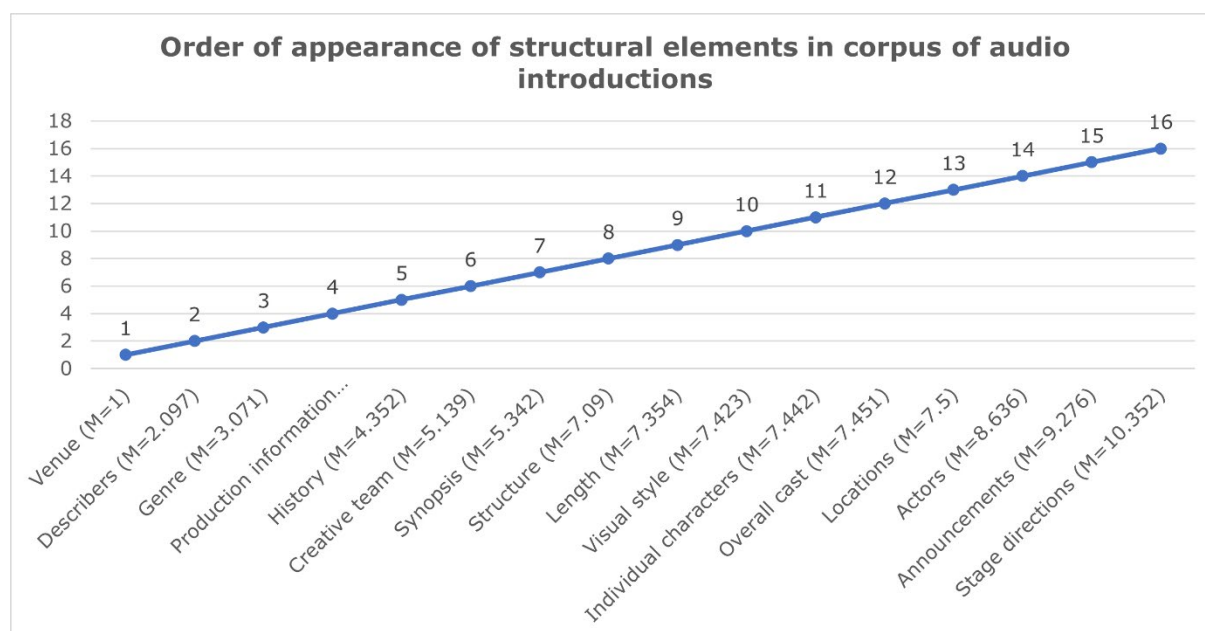


Figure 8. Order of appearance of structural elements.

In response to RQ3, our analysis shows that there is a degree of consistency in the type and order of appearance of structural elements across our corpus. We found that all major and subordinate structural elements in our framework were represented in the 52 audio introduction texts surveyed, albeit to varying degrees. This confirms prior findings in Di Giovanni (2014), Iturregui-Gallardo and Solás (2019) and Reviers *et al.* (2021) concerning the presence of broad structural themes in audio introductions, around which descriptions can be organised. Despite the lack of standardised templates for audio introductions and the flexibility this affords the genre, there is a degree of standardisation inherent in audio introduction texts, given the structural features they share. This structural consistency can serve as a basis for further alignment, in terms of, for instance, what information to prioritise and how it should be presented, depending on the specific functions the audio introduction text aims to fulfil, whether informative, narrative, expressive, persuasive, or light entertainment (Mazur 2020). These shared structural features may also facilitate the work of theatres in ensuring greater coherence and consistency across performances, as advocated by Iturregui-Gallardo and Solás (2019) and seen in our study.

However, not all structural elements were represented proportionally. We found that overall, elements pertaining to practical or extratextual information were predominant, save for the 'Individual characters' element, which was present in all 52 audio introductions. This uneven balance is perhaps related to the functional relationship between audio introductions and descriptions, with describers prioritising practical information such as credits or announcements at the outset, since intratextual information could be provided later in ADs. This functional approach may also help account for the consistency in the order of appearance of our structural elements. We note that informative structural elements such as 'Venue' or

'Announcements' tend to frame the start and end of audio introductions; this serves to signpost users throughout the performance, while including stage directions as the final element of audio introductions helps transition from the introduction to the play itself. The consistent order of appearance of structural elements across our corpus again highlights the degree of standardisation already inherent in audio introductions, which could potentially be leveraged and codified into standards of practice.

6. Conclusion

Our results concerning the linguistic complexity of audio introductions serve as evidence that describers adapt the language used in producing audio introductions to consider the audience profile and the genre of the performance. High lexical variation and complexity is avoided, to grant comprehension and access. We report a moderate LD as measured by TTR and HD-D scores for both the total corpus and the sub-corpus. Slightly more complex morpho-syntactic structures as measured by mean sentence length and FK Grade scores are observed in the total corpus versus the sub-corpus. Consistent with the VocalEyes Describing Diversity Report (Hutchinson *et al.* 2020), we identify gendered imbalances in personal characteristics descriptions, which are marked both quantitatively and qualitatively. Some of these imbalances, and particularly those relating to numerical differences by gender, are generally determined by the play and casting, rather than the describers' lexical choices. In these cases, we should consider that gendered descriptions of age, social standing, and race mirror the type of reality constructed on stage, rather than describers' choices. In other instances, imbalanced representations converge into trends of unconscious bias in the audio introduction description. In our corpus, these were most visible in descriptions of physical appearance, which were disproportionately skewed towards female characters, and value judgements, which emphasise personality and manner in men, and physical attributes in women. The decision of which traits to emphasise or efface, and *how* to construct the respective description, may therefore produce imbalances in how male and female characters are represented. While the role of the audio describer is to describe what is on stage, *how* they choose to do so may have inadvertent effects, such as proliferating gender stereotyping. We therefore believe that, in addition to studying audio introductions and descriptions, of immediate importance is the need to raise awareness of the representation and selection of characters from the point of view of theatre programme selection and access integration, as well as to embed diversity in the creative and access provision teams. Due to the local realities, our corpus included audio introductions uniquely created by female, non-disabled and white professional describers; describers of different gender, abilities or racial identities may have made different choices in constructing their descriptions, though this was beyond the scope of our study. Regarding the structural organisation of audio introductions, we find that despite the lack of predefined templates, there is a degree of standardisation inherent in audio introductions in terms of

what structural elements are present, and their order of appearance. We posit this consistency could be leveraged by individual theatres in ensuring the consistency of audio introductions across their performances, as well as codified into best practice guidelines for the wider industry.

Acknowledgements

We would like to thank the Audio Describers and the theatre Access Manager for their support and for providing access to the corpus. We are also grateful to the *JoSTrans* editing team and the reviewers for their constructive feedback and support throughout the publishing process.

References

- **Arias-Badia, Blanca and Anna Matamala** (2020). "Audio description meets Easy-to-Read and Plain Language: results from a questionnaire and a focus group in Catalonia." *Zeitschrift für Katalanistik* 33, 251–270.
- **Bardini, Floriane** (2017). "Audio Description Style and the Film Experience of Blind Spectators: Design of a Reception Study." *Rivista internazionale di tecnica della traduzione* 19, 49–73. DOI: 10.13137/2421-6763/17351.
- **Bernabé, Rocío and Pilar Orero** (2020). "Easier audio description: exploring the potential of Easy-to-Read principles in simplifying AD." Sabine Braun and Kim Starr (eds) (2020). *Innovation in Audio Description Research*. London: Routledge, 55–75. DOI: 10.4324/9781003052968-4.
- **Bernabo, Laurena** (2021). "Whitewashing diverse voices: (de)constructing race and ethnicity in Spanish-language television dubbing." *Media, Culture & Society* 43(7), 1297–1310. DOI: 10.1177/0163443721999932.
- **Bruti, Silvia and Serenella Zanotti** (2018). "Representations of Stuttering in Subtitling: A View From a Corpus of English Language Films." Irene Ranzato and Serenella Zanotti (eds) (2018). *Linguistic and Cultural Representation in Audiovisual Translation*. New York: Routledge, 228–262.
- **Chmiel, Agnieszka and Iwona Mazur** (2021). "A homogenous or heterogeneous audience? Audio description preferences of persons with congenital blindness, non-congenital blindness and low vision." *Perspectives* 30(3), 552–567. DOI: 10.1080/0907676X.2021.1913198.
- **De Marco, Marcella** (2016). "The 'engendering' approach in audiovisual translation." *Target. International Journal of Translation Studies* 28(2), 314–325. DOI: 10.1075/target.28.2.11dem.
- **De Ridder, Reglindis** (2020). "Linguistic diversity in audiovisual media for children in Belgium and Austria." Rudolf Muhr and Juan Thomas (eds) (2020). *Pluricentric Theory beyond Dominance and Non-Dominance: A Critical View*. Graz/Berlin: PCL-Press, 121–136.
- **Di Giovanni, Elena** (2014). "Audio introduction meets audio description: an Italian experiment." *inTRAlinea Special Issue: Across Screens Across Boundaries*. https://www.intralinea.org/specials/article/audio_introduction_meets_audio_description

- **Drotner, Kirsten** (2018). "Children and Media." Philip Napoli (ed.) (2018). *Mediated Communication*. Berlin/Boston: De Gruyter Mouton, 379–394. DOI: 10.1515/9783110481129-021.
- **von Flotow, Luise and Daniel E. Josephy-Hernández** (2018). "Gender in audiovisual translation studies: Advocating for gender awareness." Luis Pérez-González (ed.) (2018). *The Routledge Handbook of Audiovisual Translation. First Edition*. London: Routledge, 296–311. DOI: 10.4324/9781315717166-19.
- **Fryer, Louise** (2016). *An Introduction to Audio Description: A practical guide. First Edition*. London/New York: Routledge. DOI: 10.4324/9781315707228-13.
- **Fryer, Louise** (2018). "Staging the Audio Describer: An Exploration of Integrated Audio Description." *Disability Studies Quarterly* 38(3), <https://dsq-sds.org/index.php/dsq/article/view/6490/5093>
- **Fryer, Louise and Jonathan Freeman** (2014). "Can you feel what I'm saying? The impact of verbal information on emotion elicitation and presence in people with a visual impairment." Anna Felnhöfer and Oswald D. Kothgassner (eds) (2014). *Challenging presence: proceedings of the 15th International Conference on Presence*. Vienna: facultas.wuv, 99–107.
- **Hermosa-Ramírez, Irene** (2021). "The hierarchisation of operatic signs through the lens of audio description: A corpus study." *MonTI. Monografías de Traducción e Interpretación* (13), 184–219. DOI: 10.6035/MonTI.2021.13.06.
- **Holsanova, Jana** (2022). "A cognitive approach to audio description: production and reception processes." Christopher Taylor and Elisa Perego (eds) (2022). *The Routledge Handbook of Audio Description*. London: Routledge, 57–77. DOI: 10.4324/9781003003052-7.
- **Hutchinson, Rachel, Hannah Thompson and Matthew Cock** (2020). *Describing Diversity: an exploration of the description of human characteristics and appearance within the practice of theatre audio description*. <http://vocaleyes.co.uk/about/research/describing-diversity> (consulted 30.06.2022).
- **Iturregui-Gallardo, Gonzalo and Iris C. P. H. Solás** (2019). "A template for the audio introduction of operas: A proposal." *Hikma* 18(2), 217–235. DOI: 10.21071/hikma.v18i2.11529.
- **Kadayat, Bam and Evelyn Eika** (2020). "Impact of Sentence Length on the Readability of Web for Screen Reader Users." Margherita Antona and Constantine Stephanidis (eds) (2020). *Universal Access in Human-Computer Interaction. Design Approaches and Supporting Technologies*. Lecture Notes in Computer Science. Copenhagen: Springer International Publishing, 261–271. DOI: 10.1007/978-3-030-49282-3_18.
- **Kincaid, J. Peter, Robert Fishburne, Richard Rogers and Brad Chissom** (1975). "Derivation Of New Readability Formulas (Automated Readability Index, Fog Count And Flesch Reading Ease Formula) For Navy Enlisted Personnel." *Institute for Simulation and Training*. Naval Technical Training Command. Millington: Research Branch (RBR-8-75).
- **Mazur, Iwona and Agnieszka Chmiel** (2012). "Audio Description Made to Measure: Reflections on Interpretation in AD Based on the Pear Tree Project Data." Aline Rémuel,

Pilar Orero and Mary Carroll (eds) (2012). *Audiovisual Translation and Media Accessibility at the Crossroads*. Leiden: Brill, 173–188.

- **Mazur, Iwona** (2020). "A Functional Approach to Audio Description." *Journal of Audiovisual Translation* 3(2), 226–245. DOI: 10.47476/jat.v3i2.2020.139.
- **Mazur, Iwona** (2022). "Linguistic and Textual Aspects of Audio Description." Christopher Taylor and Elisa Perego (eds) (2022). *The Routledge Handbook of Audio Description*. London: Routledge, 93–106. DOI: 10.4324/9781003003052-9.
- **McCarthy, M. Philip** (2005). *An assessment of the range and usefulness of lexical diversity measures and the potential of the measure of textual, lexical diversity*. PhD Thesis, The University of Memphis.
- **McCarthy, M. Philip and Scott Jarvis** (2007). "vocd: A theoretical and empirical evaluation." *Language Testing* 24(4), 459–488. DOI: 10.1177/0265532207080767.
- **McCarthy, M. Philip and Scott Jarvis** (2010). "MTLD, vocd-D, and HD-D: A validation study of sophisticated approaches to lexical diversity assessment." *Behavior Research Methods* 42(2), 381–392. DOI: 10.3758/BRM.42.2.381.
- **Orero, Pilar and Sabine Braun** (2010). "Audio Description with Audio Subtitling – an emergent modality of audiovisual localisation." *Perspectives* 18(3), 173–188.
- **Orero, Pilar and Steve Wharton** (2007). "The Audio Description of a Spanish Phenomenon: Torrente 3." *The Journal of Specialised Translation* 7, 164–178.
- **Perego, Elisa** (2019). "Into the language of museum audio descriptions: a corpus-based study." *Perspectives* 27(3), 333–349. DOI: 10.1080/0907676X.2018.1544648.
- **Pérez L. Heredia, María** (2016). "Translating Gender Stereotypes: An Overview on Global Telefiction." *Altre Modernità*, 166–181. DOI: 10.13130/2035-7680/6854.
- **Ranzato, Irene** (2018). "The British Upper Classes: Phonological Fact and Screen Fiction." Irene Ranzato and Serenella Zanotti (eds) (2018). *Linguistic and Cultural Representation in Audiovisual Translation*. New York: Routledge, 203–227.
- **Remael, Aline, Nina Reviere and Gert Vercauteren** (2014). *Pictures painted in words: ADLAB audio description guidelines*. Trieste: Edizioni Università di Trieste.
- **Reviere, Nina** (2014). "Audio Introductions." Aline Remael, Nina Reviere and Gert Vercauteren (eds) (2014) *Pictures Painted in Words, ADLAB Audio Description Guidelines*. Trieste: Edizioni Università di Trieste.
- **Reviere, Nina** (2018). "Studying the language of Dutch audio description: An example of a corpus-based analysis." Laura Incalcaterra McLoughlin, Jennifer Lertola and Noa Talaván (eds) (2018). *Audiovisual translation in applied linguistics: Educational perspectives. Special issue of Translation and Translanguaging in Multilingual Contexts* 4(1). Amsterdam: John Benjamins, 178–202. DOI: 10.1075/ttmc.00009.rev.
- **Reviere, Nina and Aline Remael** (2013). "Audio Description for live performances: The role of the introduction." Paper presented at 5th *Advanced Research Seminar on Audio Description (ARSAD)* (Barcelona, 13–14 March, 2013).
- **Reviere, Nina and Hanne Roofthoof** (2018). "The functions of AIs: and up-date." Paper presented at *Languages and the Media* (Berlin, 3–5 October, 2018).

- **Reviere, Nina, Hanne Roofthoof and Aline Remael** (2021). "Translating multisemiotic texts: The case of audio introductions for the performing arts." *The Journal of Specialised Translation* 35, 69–95.
- **Romero-Fresco, Pablo** (2022). "Audio introductions." Christopher Taylor and Elisa Perego (eds) (2022). *The Routledge Handbook of Audio Description*. London: Routledge, 423–433.
- **Romero-Fresco, Pablo and Louise Fryer** (2014). "Audio introductions." Anna Matamala, Anna Maszerowska and Pilar Orero (eds) (2014). *Audio Description: New Perspectives Illustrated*. Amsterdam: John Benjamins, 11–28.
- **Snyder, Joel** (2014). *The visual made verbal: a comprehensive training manual and guide to the history and applications of audio description*. Arlington: American Council for the Blind.
- **Soler Gallego, Silvia** (2018). "Audio descriptive guides in art museums: A corpus-based semantic analysis." *Translation and Interpreting Studies* 13(2), 230–249.
- **Soler Gallego, Silvia** (2019). "Defining subjectivity in visual art audio description." *Meta: Translators' Journal* 64(3), 708–733. DOI: 10.7202/1070536ar.
- **Soler Gallego, Silvia and Olalla M. Luque** (2018). "Paintings to my ears: A method of studying subjectivity in audio description for art museums." *Linguistica Antverpiensia, New Series* 17, 140–156. DOI: 10.52034/lanstts.v17i0.472.
- **Szarkowska, Agnieszka** (2011). "Text-to-speech audio description: towards wider availability of AD." *The Journal of Specialised Translation* 15, 142–162.
- **Taylor, Christopher and Elisa Perego** (2022). *The Routledge Handbook of Audio Description*. London: Routledge. DOI: 10.4324/9781003003052.
- **Templin, Mildred** (1957). *Certain language skills in children: Their Development and Inter-relationships*. Minneapolis: University of Minnesota Press.
- **Thompson, Hannah** (2018). "Audio Description: Turning Access to Film into Cinema Art." *Disability Studies Quarterly* 38(3). <http://dsq-sds.org/article/view/6487/5085> (consulted 30.06.2022).
- **Tor-Carroggio, Irene** (2020). "T(ime) T(o) S(tart) synthesising audio description in China? Results of a reception study". *The Journal of Specialised Translation* 34, 171–191.
- **Vercauteren, Gert** (2007). "Towards a European guideline for audio description." Jorge Díaz Cintas, Pilar Orero and Aline Remael (eds) (2007). *Media for All: Subtitling for the Deaf, Audio Description and Sign Language*. Amsterdam: Rodopi, 139–150.
- **York, Greg** (2007). "Verdi made visible: audio introduction for opera and ballet." Jorge Díaz Cintas, Pilar Orero and Aline Remael (eds) (2007). *Media for All: Subtitling for the Deaf, Audio Description, and Sign Language*. Amsterdam: Rodopi, 215–230.

Websites

- **Accessible Media Inc.** (2014). *Integrated Described Video*. <https://www.ami.ca/sites/default/files/2020-07/IDV%20White%20Paper.pdf> (consulted at 30.06.2022).
- **Bansal, Shivam** (2016). *Textstat*. <https://github.com/shivam5992/textstat> (consulted 30.06.2022).
- **Hearle, Noah** (2007). *Sentence and Word Length*. <http://hearle.nahoo.net/Academic/Maths/Sentence.html> (consulted 30.06.2022).
- **Netflix** (2020). *Audio Description Style Guide v2.5*. <https://partnerhelp.netflixstudios.com/hc/en-us/articles/215510667-Audio-Description-Style-Guide-v2-3> (consulted 30.06.2022).
- **Ofcom** (2000). *ITC Guidance on standards for audio description*. http://audiodescription.co.uk/uploads/general/itcguide_sds_audio_desc_word3.pdf (consulted 30.06.2022).
- **Ofcom** (2021). *Ofcom's Guidelines on the Provision of Television Access Services*. https://www.ofcom.org.uk/__data/assets/pdf_file/0025/212776/provision-of-tv-access-services-guidelines.pdf (consulted 30.06.2022).
- **QSR International Pty Ltd.** (2020). <https://www.qsrinternational.com/nvivo-qualitative-data-analysis-software/home> (consulted 30.06.2022).
- **Shen, Y.S. Lucas** (2018). *Lexical Richness*. <https://github.com/LSYS/LexicalRichness> (consulted 30.06.2022).

Data availability statement

Data subject to third party restrictions – the data that support the findings of this study can be requested from the corresponding author, Alina Secară. Restrictions apply to the availability of these data, which were used under license for this study and third-party permission for further availability is needed.

Biographies

Alina Secară is Senior Scientist at the University of Vienna Centre for Translation Studies, where she investigates accessibility practices and technologies with a focus on cultural settings, and teaches subtitling, captioning and multimedia localization processes and technologies. She is UK Stagetext accredited theatre captioner and she worked in EU-funded DigiLing, eCoLoTrain and eCoLoMedia projects contributing to the creation of multilingual, multimedia e-learning resources for digital linguists.

E-mail: alina.secara@univie.ac.at

ORCID: 0000-0001-6281-5035



Raluca Chereji is a doctoral student and HAITrans research assistant at the Centre for Translation Studies, University of Vienna. She holds an MSc in Specialised Translation (Scientific, Technical and Medical) from University College London. Prior to her role, Raluca worked as a Project Manager and freelance medical translator. Her research interests include expert-to-lay communication, the use of automatic speech recognition in patient-facing medical translations, and narrative medicine.

E-mail: raluca-maria.chereji@univie.ac.at

ORCID: 0000-0001-6112-0856

