



Rethinking dubbing workflows: The tentative role of pre-editing in machine-translated content

Laura Mejías-Climent

Jaume I University
Castellón de la Plana, España
lmejias@uji.es 

<https://orcid.org/0000-0003-2933-7195> 

Alejandro Romero-Muñoz

Jaume I University
Castellón de la Plana, Spain
alromero@uji.es 

<https://orcid.org/0000-0002-5869-0652> 

Abstract: The DubTA project investigates the integration of neural machine translation into dubbing workflows. Given that audiovisual translation has historically been less permeable to machine translation due to its inherently creative and multimodal nature, this article first reviews the recent increase in automation technologies within audiovisual translation. This context motivated the DubTA project, whose methodology and preliminary findings are presented here. The primary goal of DubTA is to evaluate the feasibility of incorporating machine translation into dubbing processes. To this end, raw output generated by two machine translation engines was analyzed, with errors systematically categorized using a custom taxonomy to identify areas suitable for potential pre-editing in fictional dubbing scripts. Based on these findings, the project explores the possibility of developing pre-editing guidelines that could help enhance machine translation output and facilitate dubbing workflows by reducing the need for extensive post-editing. While promising, these results highlight the need for further research to refine these preliminary guidelines and assess their impact on diverse dubbing scenarios.

Keywords: audiovisual translation; dubbing; machine translation; pre-editing; post-editing.

1. Introduction

In recent years, technological advancements have led to an exponential increase in audiovisual production of all kinds and genres (Bolaños-García-Escribano et al., 2021). This surge has been met with a corresponding rise in demand all across the globe, creating a dynamic and ever-evolving market where translation and localization processes have become crucial (Rodríguez Fernández-Peña, 2023). What is more, the emergence of technological innovation and artificial intelligence (AI) has encouraged new human-computer patterns of interaction and inevitably reshaped professional landscapes (Rico Pérez & Sánchez Ramos, 2023). Among these technological advancements, the adoption of machine translation post-editing (MTPE) has become increasingly



prevalent, fundamentally altering the roles of translators and resulting in new professional profiles and educational requirements, in addition to modifying translators' attitudes and motivations (Sakamoto et al., 2024).

Machine translation (MT) systems have been available for several decades now, but, particularly since the emergence of engines based on neural networks, this technology has been organically embraced in fields that employ highly standardized terminological patterns, denotative and objective language, and even controlled language, such as specialized translation in scientific or technical fields. These neural MT (NMT) systems together with the development of AI tools are also facilitating the integration of modern technologies into more creative domains such as literary (Guerberof Arenas & Toral, 2020), audiovisual translation (AVT) (de Los Reyes Lozano & Mejías-Climent, 2023b) and even video game localisation (Anselmi & Rubio, 2020; Hansen & Houlmont, 2022). Nonetheless, this transition considerably calls into question the quality of both raw (non-edited) and post-edited output and raises alarms regarding user experiences, ethical considerations (Costa & Silva, 2020; Moniz & Parra Escartín, 2023), labor conditions (do Carmo, 2020; do Carmo & Moorkens, 2022, among others), and translators' and trainees' perceptions (Koglin et al. 2023; Calvo-Ferrer, 2024). Research into the intersection of MT and AVT is thus increasingly necessary to understand the impact of modern technologies on all stakeholders, including professionals, students, educators, researchers, audiences, and society in general.

To this respect, AVT represents an area whose multimodal nature (Kress & Van Leeuwen, 2001), combined with the wide range of genres included in audiovisual texts (Chaume, 2004), have challenged the use of MT engines (Bolaños García-Escribano, 2024). In other words, AVT involves processing information generated through a set of meaning codes (sometimes also called "signifying codes") and channels that do not only involve written text—easily processed by MT engines—but also acoustic and visual content; in short, multimodal texts (Díaz-Cintas, 2020; Remael & Reviers, 2019). In addition, the wide variety of audiovisual genres that do not adhere to the rules of a delimited terminological domain and the presence of dialogues intended to sound natural and realistic make it even more difficult to process the content employing MT engines (Mejías-Climent & de los Reyes Lozano, 2021). These factors meant that the first projects focused on combining MT and AVT were not developed until almost a decade ago, as will be detailed in the next section.

Regarding MT practices, post-editing (MTPE) is always included in professional workflows to ensure that the raw output generated by MT is acceptable considering the objectives of the project and the characteristics and context of the text that is being translated (Sánchez Ramos & Rico Pérez, 2020). However, pre-editing (PrE) does not seem to be as widely utilized as MTPE since it entails preparing the text for processing by the MT engine, which would involve an extra step in the professional workflow. Among other reasons worth exploring, this may be why the use of PrE in AVT does not seem to receive as much scholarly attention as MTPE. Consequently, this gap presents a compelling rationale to focus on PrE as part of the interests of the research project DubTA, which will be introduced in section 2.3.

In such context, this article presents the preliminary results of the analysis of MT output generated by two different engines for dubbing fictional animation material, a genre that remains largely underexplored in terms of MTPE use. The aim of this contribution is twofold: first, to employ innovative methodology in a corpus-based study to empirically analyze the use of MT in dubbing.



Second, to present the data obtained in DubTA as the first step towards designing a future set of guidelines for pre-editing the source text in AVT, a procedure that might be advantageous for streamlining purposes (Hiraoka & Yamada, 2019; Ortiz-Boix, 2016), but which seems to remain underexplored (Bolaños & Declercq, 2023). In this regard, the emerging profile of the pre-editor is further consolidated, paving the way for future studies analyzing to what extent MT optimizes dubbing workflows for audiovisual translators and generates modern human-computer interactions.

2. MTPE, creative content and the need for PrE

In an ever-progressing technological landscape, it seems imperative to analyze and systematize the changes brought about by automation. Within the realm of AVT, the integration of MTPE had been less evident until recently (Georgakopoulou, 2019), unlike other specialized domains, owing to the inherently creative and multimodal nature of audiovisual content—yet a significant portion of contemporary audiovisual content includes educational or corporate videos, where MT has shown promising results over the past five years (Wang, 2023). In any case, the inevitability of MTPE expanding to more creative fields has become apparent in the last few years (Hadley et al., 2022; O'Hagan, 2020). Among other ethical and conceptual considerations (Villanueva Jordán & Romero-Muñoz, 2023), the incorporation of MTPE into the professional sphere engenders novel interactions between humans and computers that are worth analyzing, such as the evolving profile of post-editors (de los Reyes Lozano & Mejías-Climent, 2023a). In this section, we will delve into the main research trends combining MTPE and AVT that have been followed recently, which will lead to outlining the context in which the DubTA project and its objectives arise.

2.1 A technological evolution reaching most AVT modes

All AVT modes have been closely linked to technological development since the inception of the AVT field, initially associated with the film industry and later developed with television broadcasts (Georgakopoulou, 2020). The introduction of magnetic tapes in the 1940s and the advent of personal computers in the 1980s marked significant milestones in the evolution of AVT. Digitalization in the 1990s laid the groundwork for audiovisual globalization, driving the demand for translation and localization services. During the 1990s, significant technological advancements emerged, such as the revolutionary automatic speech recognition (ASR) programs, also known as speech-to-text tools, and translation memory-based computer-assisted translation tools (Georgakopoulou, 2019). Subsequently, multilingual subtitle creation became more efficient through the implementation of subtitle templates. Later, cloud-based platforms like Oona (for subtitling) and ZooDubs (for dubbing) (Rodríguez Fernández-Peña, 2023) were incorporated into professional workflows along with recent developments in text-to-speech and speech-to-speech programs. Recently, MTPE has significantly impacted AVT workflows, with NMT emerging as the most promising branch. In addition, over-the-top services have significantly expanded audiovisual services in the last decade (Bolaños-García-Escribano et al., 2021).

In the case of MTPE, its application in the AVT field, particularly for fiction and entertainment genres, was limited until the past decade (Rico Pérez, 2023). The semiotic complexity of audiovisual



content and the need for interpretation and creativity have posed challenges to the integration of MTPE into AVT workflows. Furthermore, the diversity of genres and the lack of specialized terminology in audiovisual texts have hindered the use of MT engines in this domain. However, in recent years, there has been a growing effort to explore the integration of MT in AVT, particularly in subtitling, given the requirement of dedicated software for this process; such technology offers an excellent platform into which MT tools could be incorporated.

In particular, several authors have investigated the application of MT systems in interlingual subtitling, focusing on the quality of the raw output, the role of MTPE, and productivity improvements (Díaz Cintas & Massidda, 2020). Major European projects such as MUSA (Multilingual Subtitling of Multimedia Content, 2002-2004) and SUMAT (Subtitling by Machine Translation, 2011-2014) have explored the combination of MT and subtitling. Additionally, the Universitat Politècnica de València has developed poliTrans (2017), an innovative online service for multilingual transcription and automatic subtitling of educational audiovisual content. Using the same technology as poliTrans, the European project EMMA (European Multiple MOOC Aggregator Project, 2014-2016) focuses on transcription, MT and the subtitling of educational videos. Similarly, the TraMOOC project (Translation for Massive Open Online Courses, 2015-2018), funded by the European Commission, provides subtitling services in 11 language pairs for online educational videos. ASR technology has also been used for automatic subtitling on platforms like YouTube and in the SAVAS project (Live Subtitling and Captioning Made Easy, 2012-2014), focused on developing independent ASR technology for generating multilingual subtitles for television broadcasts. Furthermore, since 2022, broadcasts by the Spanish national TV channel La 1 de TVE include 17 regional news programs with automatic subtitles (de Higes Andino, 2023).

YouTube also introduced accessibility options in 2008 and has offered automatic subtitling since 2010, relying on Google's ASR technology and the NMT engine Google Translate. Following the model of subtitling and dedicated software incorporating MTPE, at the Universitat Autònoma de Barcelona, several projects combining MT and audio description have been conducted. One example is ALST (Linguistic and Sensorial Accessibility: Technologies for Voiceover and Audio Description, 2013-2015), which compares the process of human audio description versus MTPE. Another project is MeMAD (Methods for Managing Audiovisual Data, 2018-2020, Aalto University School of Science in Finland), aimed at developing a methodology for the efficient reuse of audiovisual content, particularly from television and over-the-top platforms.

2.2 From script to screen: MTPE meets dubbing

According to de los Reyes Lozano & Mejías-Climent (2023b), dubbing and AI have been key for companies' investments since 2022, as demonstrated by startups like Dubdub, Dubverse, and Deepdub. These companies primarily aim to automate dubbing processes, particularly for educational videos. Technological development has also made it possible to imitate or clone the original characters' voices to dub their utterances, as software programs such as Aloud, Rask.ai, Dubbah and HeyGen have demonstrated recently.

Despite these examples, research experiences using MT for dubbing purposes are somewhat less common than for subtitling, but in both cases, they tend to cover all stages, from transcription



to voice synthesis. This use is typically restricted to non-fiction genres, whose audiovisual configuration and more restricted field and terminology make the MT output more accurate. In addition, great emphasis is placed on speech-to-text, text-to-speech and speech-to-speech processes—since they are technologically complex and costly stages—yet the linguistic transfer using MT seems to be taken for granted, with little examination of the challenges and limitations it may entail. Be it as it may, four major projects specifically focus on the integration of MT into dubbing: Matousek and Vít (2012) use MT results in subtitling to adapt the translated text for dubbing purposes. Secondly, Marcello Federico's team developed software for Amazon automating the entire dubbing process, focusing particularly on voice synthesis (Federico et al., 2020; Tam et al., 2022). Thirdly, Marco Turchi's team at the Fondazione Bruno Kressler (Trento) seeks to implement dubbing automation strategies, distinguishing between on-screen and off-screen dubbing (Karakanta et al., 2021). In addition, Yang et al. (2020) presented a system designed for large-scale dubbing involving transcription, translation and voice synthesis using the original speaker's voice and recreating lip movements to synchronize with the translated audio.

As evidenced by the number of recent projects and the findings from the latest European Language Industry Survey, the integration of MTPE into professional contexts has become tangible. The role of professional translators is transforming, extending beyond traditional translation tasks to encompass the review of MT-generated output. This transition represents a broadening of translators' responsibilities. To regulate practices related to MTPE, the ISO standard 18587:2017 was introduced, delineating requirements for translation service providers, clients, and post-editors. De los Reyes Lozano and Mejías-Climent (2023a) scrutinized the similar UNE ISO standard 18587 (AENOR, 2020), implemented three years later with analogous objectives to the ISO standard, and proposed adaptations tailored to MT for dubbing. These standards underscore the need for new professional profiles capable of proficiently handling MTPE for specific purposes. Moreover, these profiles require expertise in both MT and AVT processes, coupled with strong translation competencies encompassing linguistic, cultural, technical, and research aptitudes.

2.3 Bridging the gap between PrE, MT and dubbing: objectives of DubTA

In this context, human-computer interaction is increasingly consolidating its presence within the professional dubbing market, a trend mirrored by a growing body of academic research. However, as explained above and pointed out by Flis and Senda (2022) and Valdeón (2022), dubbing appears to be the most complex AVT mode to be processed through automation, and research and professional examples dealing with PrE do not seem to be very common, as the focus has primarily been on MTPE. The former refers to the process of modifying a source text before it is processed by an MT system to enhance output quality and reduce the need for extensive post-editing (Zhang, 2025). It is commonly categorized into two methods: bilingual and monolingual PrE (Hiraoka & Yamada, 2019). Bilingual PrE involves adjusting the source text while simultaneously reviewing the MT output, ensuring that modifications effectively improve translation accuracy. In contrast, monolingual PrE is performed without access to the MT output, meaning that it does not require proficiency in the target language (Miyata & Fujita, 2021). Regardless of the approach taken, research in other fields has demonstrated that PrE can enhance MT quality by reducing ambiguity and refining



structures before translation, thereby minimizing common errors and decreasing MTPE effort (Miyata & Fujita, 2021; Ortiz-Boix, 2016; Sakaguchi et al., 2024; Zhang, 2025). As highlighted by Sakaguchi et al. (2024), PrE can contribute to producing more natural translations:

[...] A user can pre-edit the text to ensure accurate transmission of the intended meaning in the target language. [...] The example above suggests the necessity of running the cycle of MT, the evaluation of the target language output, and the modification of the source language input (Sakaguchi et al., 2024, p. 8605).

Among others, the aforementioned studies have demonstrated that PrE improves MT quality, reduces the MTPE effort required to achieve an acceptable output, and helps prevent common MT errors. However, its application in AVT remains largely unexplored, with scholarly discourse predominantly focused on MTPE. Given the unique constraints of dubbing, where synchrony and naturalness are crucial, further investigation into the feasibility and advantages of PrE in this domain is required. In this sense, the DubTA project (Machine Translation Applied to Translation Processes for Dubbing, 2021–2022) aims to fill this gap with empirical data on the results provided by two NMT engines applied to dubbing for animated content. In particular, the project tentatively explores the incorporation of PrE into dubbing translation by taking advantage of the script preparation phase (Chaume, 2012). However, before empirically testing the actual impact of PrE on post-editing effort, it is necessary to identify and systematically categorize the types of errors generated by MT in dubbing-oriented material. This paper therefore focuses on this preliminary stage, namely the compilation and analysis of recurrent MT errors in animated dubbing scripts, as a necessary step toward the subsequent development of informed PrE guidelines. In other words, as a first step toward the overall objective of the DubTA project, it will first be necessary to systematically categorize frequent errors generated by MT in dubbing so that we can begin marking them in an exploratory PrE design phase.

With regard to the professional dubbing translation process, this workflow appears to offer a natural setting for the integration of MTPE during the script adaptation and preparation phase, which could benefit from an initial PrE of the source text. The adaptation phase of the translated script involves segmenting the text into loops and applying dubbing symbols and synchronization markers (Cerezo Merchán et al., 2016; Chaume 2012; Spiteri Miggiani, 2019). As a working hypothesis, the DubTA project proposes that the script could also be post-edited appropriately during the script adaptation step, and this MTPE process would be relatively straightforward if the text had been pre-edited beforehand, prior to generating the script through MT. Given the limitations of the project, as will be outlined in the following sections, this hypothesis will not yet be tested. The preliminary set of PrE guidelines derived from the categorization of frequent errors still requires further research to refine and assess them experimentally. However, this will provide a clear starting point for the continuation of the project.

The DubTA project, conducted by researchers from the TRAMA group at Universitat Jaume I from January 2021 to December 2022, aimed to explore the streamlining of dubbing translation and script adaptation workflows through the use of NMT. The main objective was thus to design preliminary PrE guidelines for preparing and labeling original scripts to facilitate translation into multiple languages using NMT. In addition, the project aimed to achieve a series of secondary goals,



the first of which focuses on the development of an innovative labeling proposal for audiovisual texts to automate subsequent translation and segmentation tasks, including the insertion of dubbing symbols and the full preparation of the dubbing script; in other words, an adapted PrE proposal considering the multimodal nature of audiovisual products to be dubbed. The results of this secondary objective are discussed in sections 4.5 and 5.

Another secondary objective was to evaluate the efficacy of various NMT programs in AVT, particularly in dubbing. To such an end, DeepL Translator and Amazon Translate (incorporated into the subtitling platform Oona) were chosen based on availability and funding considerations (Villanueva Jordán & Romero-Muñoz, 2023). More specifically, these tools were selected as they were the most accessible solutions at the time of the project's implementation, both in economic terms and in line with expert recommendations within the audiovisual translation sector. In addition, DubTA also aimed to establish a metric for time and effort to assess the use of NMT in dubbing and determine problematic areas that could potentially be pre-edited (see section 3 and Villanueva Jordán & Romero-Muñoz, 2023). By pursuing these objectives, the DubTA project sought to provide empirical insights into NMT's integration into dubbing. The methodology of this project was designed to gather both qualitative and quantitative data, thereby furnishing valuable insights into the prospective advantages and challenges associated with NMT in dubbing workflows. While the project primarily focused on English to Castilian Spanish translation, the optimized PrE approach could potentially accrue benefits for other target languages, similar to subtitling templates. Through this comprehensive approach, the DubTA project aimed to contribute to the advancement of NMT technology in the field of AVT.

3. Methodology

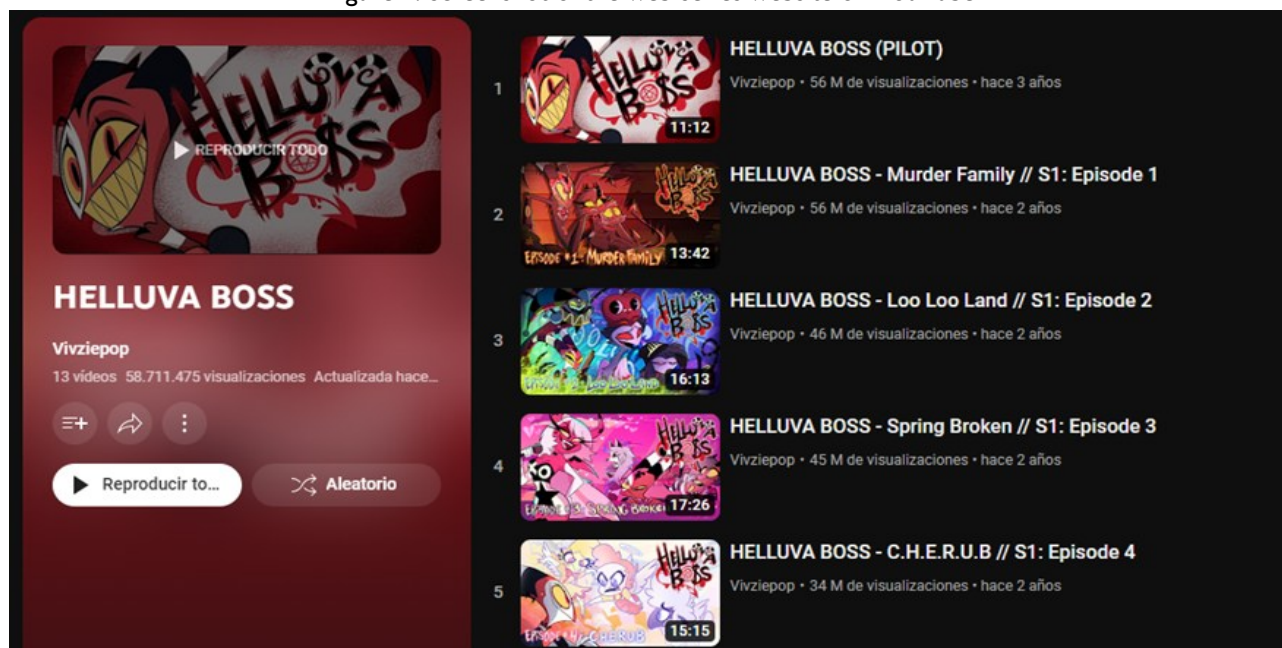
The starting point of the study was the selection of the corpus and the research tools. With regard to the latter, two MT engines, Amazon Translate and DeepL Translator, were selected to analyze how they translated an audiovisual corpus from the animated series *Helluva Boss* (Vivziepop, n.d.) from English into Castilian Spanish. To select the corpus, a comprehensive search for fictional animated videos available online was conducted. Analyzing animated material was considered, firstly, due to the limited presence of fictional corpora in empirical studies (see sections 2.1, 2.2), which appear to pose greater processing challenges for MT engines. This is largely attributed to the high level of spontaneity and naturalness in the discourse, as well as the absence of controlled or specialized language. The presence of dialogues intended to sound natural, spontaneous, and even vulgar, including humor and sarcasm, which makes it more challenging to be processed using MT. Additionally, professional demands regarding dubbing synchronies in such productions are not as imperative as they are in other fictional audiovisual genres involving real-life actors (Chaume, 2012). This makes animation a good starting point for testing MT output for dubbing purposes and developing a consequent preliminary proposal for PrE. At the same time, it should be noted that the corpus is limited to this specific audiovisual genre. The findings derived from this material may therefore not be directly generalizable to live-action productions or to other genres where technical constraints play a more decisive role. As such, the present study should be understood as an



exploratory case study that provides a starting point for further research across a broader range of audiovisual formats.

Considering these characteristics and the availability of fiction material online, the web series *Helluva Boss* was selected, a black comedy series consisting of two seasons streamed on YouTube. When this product was chosen, *Helluva Boss* was available only in English and five of the six season-one episodes had aired. These five episodes (the pilot and episodes 1 through 4), transcribed and translated, represent the bilingual and written corpus of analysis for this project. The episodes vary in length, between 11 and 17 minutes, as shown in Figure 1.

Figure 1: Screenshot of the web series website on YouTube



Source: HELLUVA BOSS channel on Youtube (Vivziepop, n.d.)

[Description] The image is a screenshot from the YouTube channel *Helluva Boss* showing, centered, a list of available episodes. Each episode is displayed with a thumbnail, accompanied on the right by the episode title, author (Vivziepop), duration, and release date. In this screenshot, the episode list includes the following: “Pilot” (season 1, 11:12 minutes), “Murder Family” (season 1, episode 1, 13:42 minutes), “Loo Loo Land” (season 1, episode 2, 16:13 minutes), “Spring Broken” (season 1, episode 3, 17:26 minutes), and “C.H.E.R.U.B.” (season 1, episode 4, 15:15 minutes) [End of description].

The second aim of this paper is to present the DubTA project as a first step towards integrating NMT into the dubbing workflow through the figure of the pre-editor. To do so, having selected this corpus, text preparation and text analysis phases were designed and implemented to trace the problematic areas in the English source text that resulted in translation errors in the Spanish versions generated by both MT engines, and which could require some form of PrE. These problematic areas were dealt with considering the specificities of dubbing (Chaume, 2012; Spiteri Miggiani, 2019).

Regarding the text preparation phase, the corpus dialogues were transcribed using Trint. With these transcripts as the source text, they were revised to verify that there were no unnecessary line breaks or punctuation marks (such as periods or ellipses); otherwise, the MT engines would divide the text into nonsensical segments. As a result, 734 segments were obtained,

using each complete character intervention as the segmentation criterion. Once the source texts were prepared, they were translated by the Amazon Translate and DeepL MT engines into Castilian Spanish (the translated text). Afterward, source text and target text segments were aligned in double-entry tables in Microsoft Word documents, and they were then turned into PDF documents (Figure 2), representing our bilingual corpus, to ensure that the alignment was maintained in the coding software during the text analysis phase. Thus, a total of 1,438 aligned segments were obtained across both MT engines, with 734 segments per engine.

Figure 2: Sample of transcription (left) and MT raw output (right)

1	HELLUVA BOSS - Loo Loo Land __ S1_ Episode 2.mp4	HELLUVA BOSS - Loo Loo Land __ S1_ Episodio 2.mp4
2	Octavia [00:00:06] Mommy! Daddy!	Octavia [00:00:06] ¡Mamá! ¡Papá!
3	Stolas [00:00:10] Via is calling, Stella.	Stolas [00:00:10] Via está llamando, Stella.
4	Stella [00:00:13] You get up.	Stella [00:00:13] Te levantas.
5	Stolas [00:00:18] Via, what troubles you, my owlet?	Stolas [00:00:18] Via, ¿qué te preocupa, mi lechuza?
6	Octavia [00:00:20] Daddy, daddy, I had a dream, a really bad dream.	Octavia [00:00:20] Papá, papá, he tenido un sueño, un sueño muy malo.
7	Stolas [00:00:28] A nightmare.	Stolas [00:00:28] Una pesadilla.
8	Octavia [00:00:30] I was looking all over the palace and I couldn't find you anywhere. You weren't there.	Octavia [00:00:30] Estuve buscando por todo el palacio y no te encontré por ningún lado. No estabas allí.

Source: Authors (2026)

[Description] The image displays a table with three columns. The first column lists the segments, numbered from 1 to 8. The second column contains segments from the source text in English, including the character's name, the timecode, and the cue. The third column presents the same information translated into Spanish (target text) [End of description].

After that, in the text analysis phase, two projects were created in the qualitative analysis software Atlas.ti. One project focused on the five-episode corpus in English and the Spanish translation by DeepL, while the other project proceeded in the same way with the Amazon Translate raw output. Each project was led by a coder, i.e., two researchers dealt with the textual material. The episodes were analyzed in iterative stages and internal validation was ensured to achieve consistency in the code application, so decisions had to be evaluated to guarantee the reliability of the data and results. For validation purposes, a qualitative approach was adopted. The two coders (one of whom is the co-author of this paper) conducted a cross-review process, systematically examining and discussing each other’s annotations to ensure consistency in the application of the coding scheme and to reduce as much as possible any biased assessment.

It should be noted that this coding process stemmed from an initial code classification (Table I), which, in turn, resulted from the combination of several quality assessment criteria (Mejías-Climent & Villanueva-Jordán, 2023), namely, the typology of errors from the Multidimensional Quality Metrics (Lommel, 2018; Lommel et al., 2014), the Dynamic Quality Framework (Görög, 2014), and the model on dubbing quality proposed by Spiteri Miggiani (2022) (Villanueva Jordán & Romero-Muñoz, 2023). Accordingly, the classification employed was specifically

designed within the DubTA project to align with the particularities of MT used in dubbing, integrating established error detection taxonomies from both MT and dubbing studies. Drawing on this initial code classification, coders progressively identified a series of patterns indicating how errors manifest in our corpus, which could ultimately serve as labels for the PrE to mark the problematic areas in MT applied to dubbing.

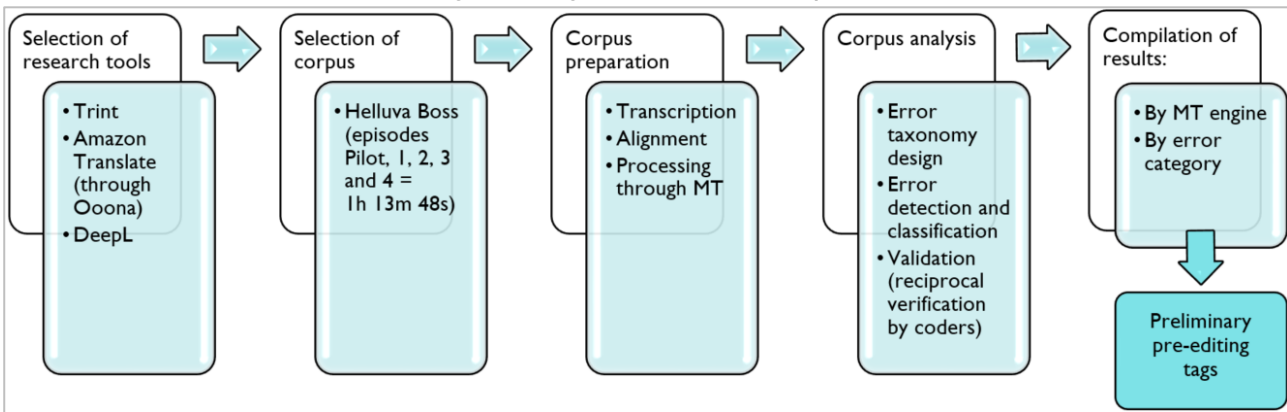
Table I: Initial code classification

Category	Code
Comprehension	COMP
Coherence/consistency	COHE
Expression	EXPR
Pragmatic-stylistic	PRAG
Technical	TECH

Source: Authors (2026)

According to the combination of quality assessment criteria, the possible errors could belong to any of the following categories: comprehension, coherence/consistency, expression, pragmatic-stylistic, and technical. Comprehension errors would include false meanings derived from a misunderstanding of the ST that materialized as inadequate decisions in the TT (lexicon and syntax, among other language levels). Coherence errors would imply inconsistencies in the translation decisions. Expression errors would usually appear as agreement and orthotypographic errors, or an incorrect use of verb tenses, among others. Pragmatic-stylistic errors would represent an unidiomatic use of the language, unnatural constructions, inadequate use of register, dialectal inconsistencies, etc. Finally, technical errors would comprise lip or phonetic synchrony problems, or excessively long translations, for instance. This initial classification was broad enough so that further patterns within the codes could be identified based on the empirical corpus data. Regarding Table I, it is worth mentioning that the comprehension category was ultimately not considered in the analysis under the assumption that MT does not include any interpretative or cognitive processing between ST and TT. These stages described above are summarized in Figure 3.

Figure 3: Stages in the DubTA Analysis



Source: Authors (2026)

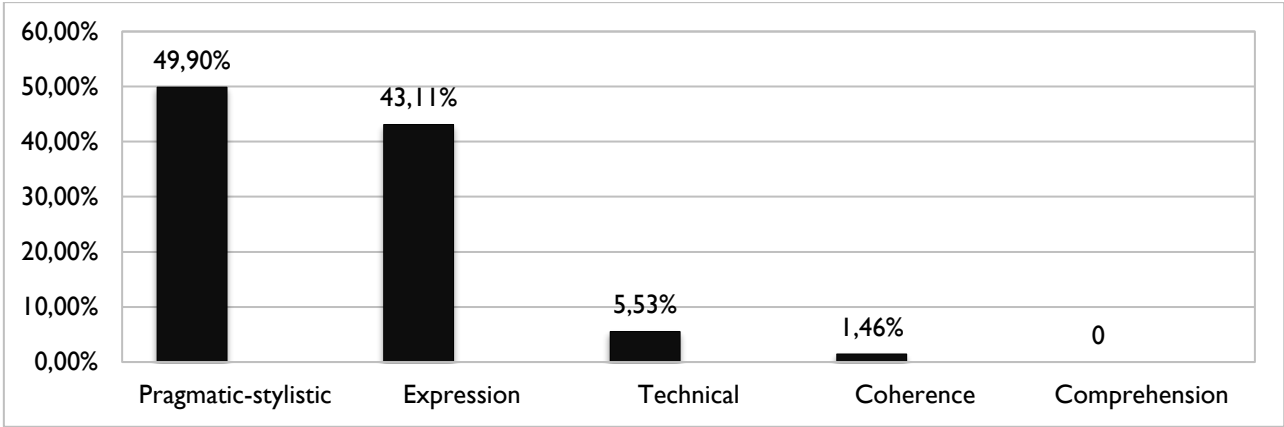
[Description] The image presents a diagram consisting of contiguous boxes connected by arrows progressing to the right, illustrating the different stages of work included in the DubTA project, namely: Selection of research tools (Trint, Amazon Translate through Ooona and DeepL), Selection of corpus (*Helluva Boss*, 1h, 13m, 48s), Corpus preparation

(transcription, alignment, processing through MT), Corpus analysis (error taxonomy design, error detection and classification, validation—reciprocal verification by coders), and finally, Compilation of results (by MT engine, by error category), leading to preliminary PrE tags [End of description].

4. Results and discussion

The corpus text analysis phase yielded results that point to possible regularities regarding the frequency of errors in our corpus, where the pragmatic-stylistic and expression categories seem to prevail. This contrasts with the rare presence of technical and coherence/consistency errors. As explained above, no comprehension errors were considered due to the nature of MT. More specifically, in our corpus comprising the five-episode MT raw output from *Helluva Boss* provided by Amazon Translate and DeepL, 958 errors were identified, 478 of which were pragmatic-stylistic (49.90%), 413 were expression errors (43.11%), 53 were technical errors (5.53%), 14 were coherence errors (1.46%), and no comprehension errors were found (Figure 4).

Figure 4: Total of errors by Amazon Translate and DeepL Translator in episodes 0 (Pilot) through 4

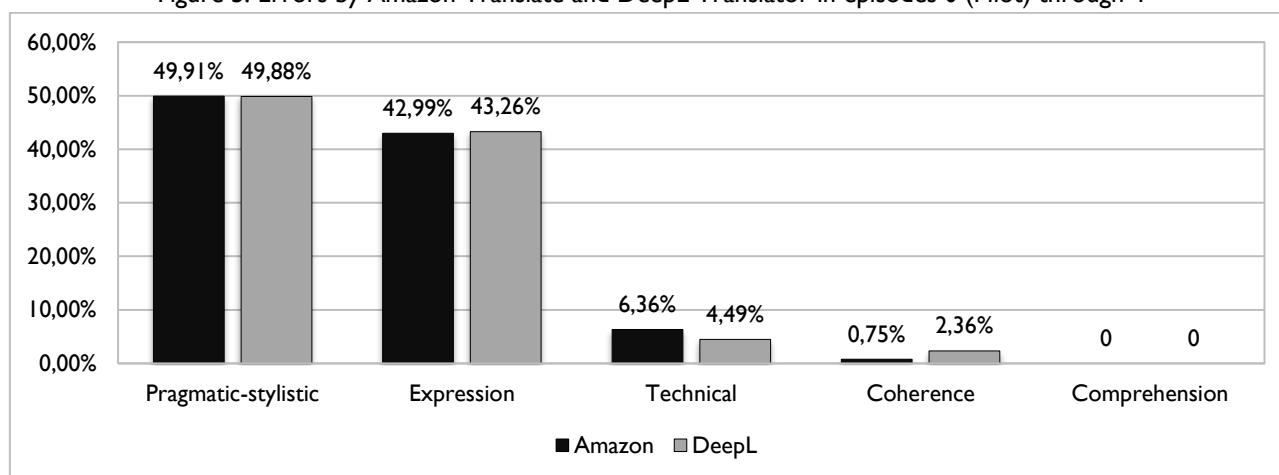


Source: Authors (2026)

Regarding the specific case of MT provided by Amazon Translate (Figure 5), 535 errors were identified. The most recurrent category was pragmatic-stylistic, with 267 errors (49.91%), followed by the expression category with 230 errors (42.99%). In contrast, the technical and consistency/coherence categories exhibited significantly lower frequencies, with 34 (6.36%) and 4 (0.75%) errors, respectively. Finally, no comprehension errors were considered. As for the MT provided by DeepL (Figure 5), 423 errors were categorized. The most recurrent category was pragmatic-stylistic, with 211 errors (49.88%), followed by the expression category, with 183 errors (43.26%). Conversely, the other categories appeared less frequently: technical errors (19 examples, 4.49%) and consistency/coherence errors (10 examples, 2.36%). Finally, no comprehension errors were identified.



Figure 5: Errors by Amazon Translate and DeepL Translator in episodes 0 (Pilot) through 4



Source: Authors (2026)

When comparing the overall results obtained by each MT engine in terms of error count, it is noteworthy that DeepL, a freely available MT engine, produced fewer errors than Amazon Translate, which operates under a paid subscription model. While this observation is based solely on the analyzed dataset and does not constitute a comprehensive quality assessment, it suggests that, within the scope of this study, the paid engine did not necessarily outperform the free one in processing fictional animated material.

With this data in mind, the following subsections will be devoted to explaining the main results obtained by the MT engines in each category.

4.1 Pragmatic-stylistic errors

According to our data, pragmatic-stylistic errors were the most common in our MT corpus. Stemming from the initial code classification, Table 2 presents the precise pragmatic-stylistic error patterns that derive from the errors found throughout the analysis.

Table 2: Pragmatic-stylistic criteria

Pragmatic-stylistic patterns	
•	Unidiomatic use of interjections, discourse markers and appellatives in Spanish.
•	Dialect. Non-systematic use of Castilian Spanish. Presence of different linguistic varieties.
•	Register. Inappropriate use of politeness markers.

Source: Authors (2026)

Focusing on the specific case of Amazon Translate, this was the most frequent error category in this MT engine, with 267 errors (49.91%) found throughout the five episodes. The most common patterns identified were classified as dialectal inconsistencies, register inconsistencies and an unidiomatic use of Spanish. Regarding dialectal inconsistencies, a systematic mixture of Castilian Spanish and other dialectal varieties was found. Although Amazon's Spanish variety was Castilian overall, sometimes the translation employed grammatical and lexical forms (e.g. *carajo*) belonging to other varieties. Regarding register, on some occasions, there was an inappropriate use of the *tuteo* (the more informal form of address associated with family, friendly, or colloquial conversations) in

situations that required more polite forms (the *usted* form of address), or vice versa (*recuerda* instead of *recuerde*, Table 3). Finally, some fragments presented an unidiomatic use of the language, particularly when it came to interjections (*maldita sea* instead of more natural options, such as *joder*), discourse markers, and appellatives.

Table 3: Inadequate register		
23	Girl 2 [00:01:18] Wait, Miss Mayberry. Remember what you taught us, think before you act.	Chica 2 [00:01:18] Espere, señorita Mayberry. Recuerda lo que nos enseñaste, piensa antes actúas.

Source: Authors (2026)

Regarding DeepL, pragmatic-stylistic errors were the most frequent category (49.91%). Similarly to Amazon Translate, the most recurrent errors could be classified as dialectal inconsistencies, register inconsistencies and an unidiomatic use of Spanish. Beyond the *tú-usted* alternation, register inconsistencies were observed with Blitzo, a character who tends to make abrupt changes in register as part of the character construction (Mejías-Climent & Villanueva-Jordán, 2023), but these shifts were not always properly translated. On the other hand, it is worth mentioning that, due to the nature of the series, many interjections used in English had a strong sexual, pain or surprise component, but they were translated into Spanish in an unidiomatic way (*oh, sí*, Table 4).

Table 4: Unidiomatic Interjections		
18	Gerald [00:01:02] OK. Oh yeah.	Gerald [00:01:02] OK. Oh, sí.
19	Martha [00:01:03] I told you we were not using this.	Martha [00:01:03] Te dije que no íbamos a usar esto.

Source: Authors (2026)

4.2 Expression errors

Expression errors were the second most frequent category found in the corpus (43.11%). Although this category encompasses various types of grammatical errors, instances could be classified as structure calques, morphological errors, grammatical interferences, syntactic errors, lexical calques and orthotypographic errors (Table 5).

Table 5: Expression criteria	
Expression patterns	
	• Structure calques: personal and possessive pronouns.
	• Morphological errors: verb tenses, agreement, grammatical category, etc.
	• Grammatical interferences: <i>ser</i> and <i>estar</i> (“to be”).
	• Syntactic errors: that clauses.
	• Lexical calques.
	• Orthotypographic errors: question marks or exclamation marks.

Source: Authors (2026)

Expression errors appeared frequently throughout the five episodes processed by Amazon Translate (42.99%). Drawing on the initial code classification, we grouped the error patterns as

structure calques, morphological errors, syntactic errors, lexical calques and orthotypographic errors. Regarding structure calques, an inadequate use of personal and possessive pronouns was found repeatedly. As for morphological errors, the wrong use of verb tenses can be mentioned, as well as errors in gender (Table 6) and number agreement. As Table 6 shows, the Spanish masculine adjective *maldito* (instead of the feminine *maldita*) did not match the grammatical gender of the feminine noun *perra*.

Table 6: Incorrect gender agreement

4	Blitzo [00:04:54] When you set fire to my office in front of a client, you fucking bitch.	Blitzo [00:04:54] Cuando prendiste fuego a mi oficina frente a un cliente, maldito perra.
---	---	---

Source: Authors (2026)

As for syntactic errors, it was problematic when the original English text omitted that within a that clause since the MT engine tended not to detect the omission. On the other hand, the most prominent grammatical interference was the confusion between *ser* and *estar* as a translation of the verb “to be” (Table 7). Finally, the most frequent spelling errors were due to the lack of question or exclamation marks, since Spanish requires both opening and ending marks (*¿*, *?*, *¡*, and *!*).

Table 7: Confusion between *ser* and *estar*

7	Ms. Mayberry [00:00:02] I was a good person before it all went down. I was good my entire life.	Sra. Mayberry [00:00:02] Fui una buena persona antes de que todo se derrumbara. Yo estaba bueno toda mi vida.
---	---	---

Source: Authors (2026)

Expression errors were also frequent (43.26%) in DeepL, where patterns in morphological errors, syntactic errors and lexical calques were noticeable. Among these, issues with gender and number agreement were the most prominent, such as the case of the word “hero”, which was translated as the masculine *héroe* (instead of the feminine *heroína*) without considering that the character is a woman. However, it should be noted that there was correct grammatical gender marking when the original text included the personal pronoun.

Table 8: Gender agreement error

43	Man 2 [00:03:19] You're a hero.	Man 2 [00:03:19] Eres un héroe.
44	Children [00:03:20] You're a hero.	Niños [00:03:20] Eres un héroe.
45	Man 3 [00:03:20] Oh, oh! You're a hero!	Man 3 [00:03:20] ¡Oh, oh! Eres un héroe.
46	Ms. Mayberry [00:03:20] She's not a hero!	Sra. Mayberry [00:03:20] ¡No es una heroína!

Source: Authors (2026)

4.3 Technical errors

The technical error category appeared with moderate frequency in the corpus translated by both Amazon Translate and DeepL (5.53%). Stemming from the initial classification, the most recurrent technical error patterns were due to synchrony problems, particularly due to the lack of isochrony in excessively long translations as well as phonetic or lip-synchrony errors (Table 9).

Regarding kinetic synchrony, no particularly salient mismatches were systematically observed in the corpus analyzed. In general terms, the tone and content of the MT output appeared to be compatible with the expressivity of the animated characters. However, this observation should be interpreted with caution. Given that the present study focuses exclusively on animated material, where gestural realism and performance constraints may differ from live-action productions, kinetic synchrony was not operationalized as an independent analytical parameter. Future research across other audiovisual genres should explicitly incorporate this dimension in order to assess how MT output interacts with embodied performance in contexts where kinetic alignment may be more decisive.

Table 9: Technical patterns

Technical patterns
• Synchronies. Lack of isochrony due to excessively long translations. Phonetic or lip-synchrony errors.

Source: Authors (2026)

Among the technical errors in Amazon Translate (6.36%), the most common issues arose because of problems with synchronies, especially isochrony and phonetic or lip-synchrony errors. Thus, the translation provided by Amazon Translate was sometimes too long to be properly voiced by dubbing actors. On the other hand, problems related to poor phonetic or lip-synchrony were relevant since the characters’ lip movements restricted the translation in certain shots.

Regarding the technical category in DeepL, accounting for 4.49% of the errors, there were patterns related to synchrony problems, particularly due to poor isochrony (as in those cases where the Spanish version was much longer than the original) or when there were phonetic or lip-synchrony problems (Table 10). In the latter case, Figure 6 shows how Martha articulates the words “thank” (shot A) and “you” (shot B), but the Spanish translation *gracias* cannot match this mouth articulation.

Figure 6: Shots of Martha saying “Oh, thank you!”



Source: *Helluva Boss* (Vivziepop, n.d.), episode I, season I (02:52)

[Description] The image shows a close-up of the character, a woman lying in bed, wrapped in bandages, with only her eyes and mouth visible. In the screenshot on the left, her mouth is open wide to produce an open vowel sound (“thank...”), while in the screenshot on the right, her mouth is shaped to form a closed vowel sound (“...you”) [End of description].

Table 10: Lip-sync error

36	Martha [00:02:51] Oh, thank you!	Martha [00:02:51] ¡Oh, gracias!
----	----------------------------------	---------------------------------

Source: Authors (2026)

4.4 Coherence or consistency errors

There was a much-reduced percentage of coherence/consistency errors within the Amazon Translate and DeepL corpus, accounting for only 1.46% of errors. Table 11 shows the basic consistency patterns: errors due to misspelled or mistranslated character names and errors due to lack of coherence between audio and image.

Table 11: Coherence or consistency criteria

Coherence or consistency criteria	
•	Misspelled or mistranslated character names.
•	Lack of coherence between visual elements (image) and sound (dialogue).

Source: Authors (2026)

In the case of Amazon Translate, only 4 coherence/consistency errors were found (0.75%). On the one hand, an occasional mistranslation of a character's name was identified (Table 12). On the other hand, there was a moment of semiotic inconsistency between the on-screen text and the information provided by the dialogue.

Table 12: Character name mistranslation

109	Millie [00:10:42] Blitz!	Millie [00:10:42] Bombing!
-----	--------------------------	----------------------------

Source: Authors (2026)

As for DeepL, 10 consistency and coherence errors were coded (2.36%). These errors can also be classified as misspelled or mistranslated character names as well as coherence problems between the image and sound. Table 13 shows a segment in which the phrase “Good Morning!” appears as on-screen text while Mrs. Mayberry simultaneously reads it aloud. In the dubbed Spanish version, the audio renders the line as *¡Buenos días!*, while the written text on screen remains in English. This creates a multimodal inconsistency, as the character is portrayed as an English teacher, yet the spoken line no longer corresponds to the instructional content displayed visually (Figure 7). The visual element cannot be modified without altering the fact that she is teaching English, which restricts the range of post-editing solutions. Although human translators would face a similar constraint, this case is particularly relevant for PrE, as it illustrates how it could anticipate multimodal conflicts. For instance, the pre-editor might propose adapting the character's role (for example, redefining her as a foreign-language teacher) or explicitly flagging the segment as requiring creative intervention at the post-editing stage.

Figure 7: Mrs. Mayberry writing on the blackboard



Source: *Helluva Boss* (Vivziepop, n.d.), episode I, season I (02:22)

[Description] The image shows a teacher holding a piece of chalk. In the left screenshot, a close-up shot captures the teacher's handwriting "Good Morning!" on the blackboard. In the right screenshot, the teacher is facing her students, with the blackboard behind her displaying the words "Good Morning!" she has just written [End of description].

Table 13: Semiotic incoherence between the image and the sound

3	Ms. Mayberry [00:00:13] Good morning! I hope you all did your homework.	Sra. Mayberry [00:00:13] ¡Buenos días! Espero que todos hayan hecho su tarea.
---	---	---

Source: Authors (2026)

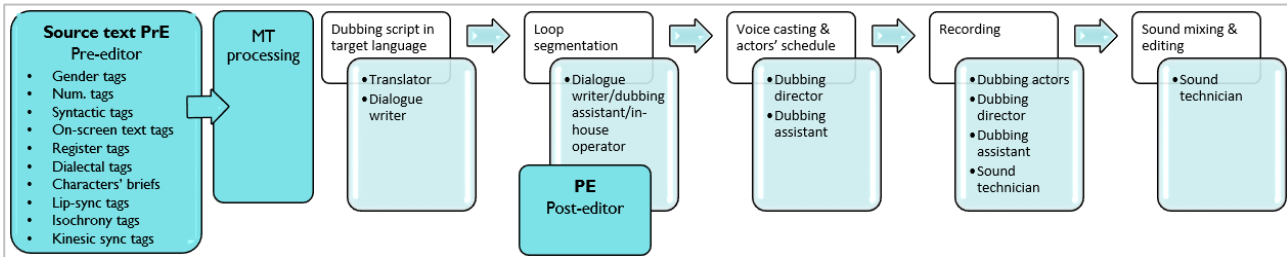
4.5 The role of the pre-editor

The aforementioned errors could be envisioned as a set of proposals to be implemented by the pre-editor in MT applied to dubbing, a role that should require considerable knowledge of both AVT and MT to ensure their successful integration. Given the evolving professional landscape, shaped by technological advancements and the diversification of roles, PrE is one of the possible tasks involved in AVT in the future. In other words, AVT translators might need to pre-edit, translate, or post-edit audiovisual texts. Further studies using a wider corpus and varied audiovisual genres would be useful to empirically validate the relevance of PrE in dubbing workflows, for which the present findings may serve as a starting point. If the corpus is enlarged and the above trends are confirmed, the largest group of errors that could be anticipated through a standardized labeling system provided by the pre-editor are the pragmatic-stylistic and expression categories. Ultimately, in line with the observations of Bywood et al. (2017), MT in general appears to yield better results in more linguistically controlled and specialized fields. However, as argued by Mejías-Climent and Villanueva Jordán (2023), it also tends to eliminate features of orality and spontaneity. This reinforces the need for the pre-editor to maintain these characteristics in creative contexts such as dubbing.

Although there have been previous experiments with PrE, it does not seem to be widely implemented, but PrE could be a positive resource for optimizing dubbing processes using NMT (Ortiz-Boix, 2016; Zhang, 2025). Despite the extent of automation that NMT offers, MTPE is always required when generating dubbing scripts that meet dubbing standards (Chaume, 2007), and this process could be streamlined by conveniently preparing the source text to be processed by a MT engine, anticipating the most frequent errors that MT tends to commit. Furthermore, translation for dubbing necessarily encompasses a rough translation phase and an adaptation phase (Chaume 2012; Spiteri Miggiani, 2019) to prepare the text and to format it as a dubbing script, so this

adaptation phase might serve as an ideal context to incorporate MTPE tasks, which, in turn, could be facilitated by PrE or the prior labeling of the source script. The following diagram presents a proposed adaptation of the dubbing process that incorporates both MTPE and PrE stages.

Figure 8: The dubbing process and main professional roles including PrE and MTPE



Source: Adapted from Spiteri Miggiani (2019, p. 7) and Chaume (2012)

[Description] The image displays a diagram of the dubbing process by stages, indicating the agents responsible for each: dubbing script in TL, loop segmentation, voice casting and actors’ schedule, recording and sound mixing and editing. The additional stages proposed in this study—PrE, processing through MT, and MTPE—are highlighted, with the latter positioned at the level of loop segmentation, and the former (PrE and MT), at the beginning of the dubbing workflow [End of description].

Among the possible PrE tasks, and considering the different error patterns identified throughout our corpus, the pre-editor could perhaps consider gender and number markers in cases where these (like personal pronouns) are absent, as expression errors indicate. The pre-editor could also provide syntactic elements, such as subject markers in incomplete sentences with ellipses or interruptions. Marking long sentences or dividing them into shorter sentences could also ease the MT process. As far as coherence or consistency proposals, the pre-editor could include tags when a problematic text appears on the screen. As for pragmatic-stylistic proposals, the pre-editor could consider dialectal markers for a homogeneous use of any given language variety; register markers for situations implying a politeness shift; the inclusion of briefs with the description of characters, etc. Finally, based on the technical category, pre-editors could include marks to ensure proper lip synchrony when shots require it, as well as include isochrony marks in extended fragments. To effectively perform these tasks, a pre-editor would need a strong understanding of dubbing conventions, familiarity with MT limitations and error patterns, and the ability to anticipate linguistic and technical issues that may arise in the MTPE process. Furthermore, a meticulous command of source language structures is essential, as many PrE decisions rely on accurately identifying and adjusting elements that could otherwise hinder MT performance.

Table 14: Elements that may need to be marked during text pre-editing according to problematic MT categories

Expression	Coherence	Pragmatic-stylistic	Technical
Gender tags	Tagging problematic on-screen text	Dialectal tags	Lip-sync tags
Number tags		Register tags	Isochrony tags
Syntactic tags		Character briefs	Kinesic sync tags
Long sentence marking or sentence shortening			

Source: Authors (2026)

Considering the data presented above, and contrary to how it might seem, the scarcity of technical or cohesion errors does not imply high quality in these categories. Conversely, the highest number of errors occurs in discourse dimensions, such as the inclusion of unnatural expressions and the inappropriate pragmatic-stylistic use of the Spanish language, which makes it infeasible to analyze other categories like technical or cohesion errors because the segments are already incorrect or inadequate. However, as mentioned, the PrE proposal outlined here is not definitive but rather a preliminary attempt to identify the categories that could benefit from a PrE phase. This tentative approach serves as a starting point for refining both the PrE process in dubbing—similarly to previous experiences with PrE and subtitling (Ortiz-Boix, 2016)—and the role of the pre-editor.

All these proposals will benefit from an extended corpus that would make it possible to increase the number of errors obtained and the examples to be considered according to each category. All in all, this tentative PrE design aims to allow MT engines to provide more accurate translations that would need a less thorough MTPE for dubbing, which would eventually make the dubbing workflow faster and easier. Nevertheless, given the exploratory scope of the present study, further research is required to ascertain the impact of PrE on the nature and incidence of errors generated by MT engines.

5. Conclusions

The first aim of this paper was to employ the methodology designed in DubTA to analyze MT output generated by DeepL Translator and Amazon Translate for dubbing purposes, focusing on a bilingual English-Spanish fictional animation corpus. DubTA's methodology followed a two-phase procedure including the text preparation (which uses a code classification combining several MT and dubbing quality assessment criteria) and the text analysis. The results of this analysis connect with our second aim since this data is envisioned as labels, which are considered a first step towards designing a proposal for PrE the ST in a dubbing context. The preliminary recommendations for PrE original scripts discussed in this article may serve as a starting point for optimizing dubbing processes incorporating the potential of NMT in creative and multimodal contexts. This PrE proposal could be refined in future research with live-action (not only animation) fiction programs.

Regarding the empirical analysis implemented within the DubTA project, the results point to a possible regularity in the frequency of errors found in both MT engines when processing our five-episode corpus from the *Helluva Boss* series. This regularity can be summarized as a preponderance of pragmatic-stylistic and expression errors, which contrasts with a very low presence of technical and coherence/consistency errors. However, these results should be carefully interpreted. The rare frequency of technical or cohesion errors does not imply a higher quality in these categories. Conversely, most errors occur as a result of an incorrect (expressive) and inadequate (pragmatic-stylistic) use of Spanish, which makes a further analysis of technical or coherence issues virtually unfeasible, because segments are already incorrect or inadequate from the beginning. Having that in mind, it seems more convenient to focus on these two major sources of linguistic inaccuracies.

Both error categories are limitations of MT. However, NMT will certainly improve over time, hence the importance of noting which dimensions of the audiovisual text need to be addressed in a potential PrE phase that could contribute to optimizing the MTPE processes. These dimensions



can in turn be mapped onto subtitling or dubbing workflows and finally operationalized into specific PrE patterns, which is our ultimate proposal for future research. Expression and pragmatic-stylistic errors may be caused by the “lack” of contextual data, which may stem from the absence of an interpretative process in MT in a multimodal context. Thus, our proposal emphasizes the need for the pre-editor as a new profile that would understand the problematic areas in the AVT workflow, particularly dubbing in this case, and would fill the needs for the AVT text to be properly translated. If MT is to make inroads into the AVT workflow, the complexity of the audiovisual text makes it necessary to think not only about quality control tasks on the final product but also about the product that has not yet been translated.

References

- Anselmi, C., & Rubio, I. (2020). The Future is Here: Neural Machine Translation for Games. *Multilingual*, 31(2). <https://multilingual.com/issues/mar-apr-2020/the-future-is-here/>
- Asociación Española de Normalización y Certificación – AENOR. (2020). UNE-ISO 18587. *Servicios de traducción. Posedición del resultado de una traducción automática. Requisitos*. UNE.
- Bolaños García-Escribano, A. (2024). *Practices, Education and Technology in Audiovisual Translation*. Routledge. <https://doi.org/10.4324/9781003367598>
- Bolaños García-Escribano, A., & Declercq, C. (2023). Editing in Audiovisual Translation (Subtitling). In C. Sin-Wai (Ed.), *Routledge Encyclopedia of Translation Technology* (2^a ed.; pp. 565–581). Routledge.
- Bolaños-García-Escribano, A., Díaz-Cintas, J., & Massidda, S. (2021). Latest Advancements in Audiovisual Translation Education. *The Interpreter and Translator Trainer*, 15(1), 1–12. <https://doi.org/10.1080/1750399X.2021.1880308>
- Bywood, L., Georgakopoulou, P., & Etchegoyhen, T. (2017). Embracing the Threat: Machine Translation as a Solution for Subtitling. *Perspectives*, 25(3), 492–508. <https://doi.org/10.1080/0907676X.2017.1291695>
- Calvo-Ferrer, J. R. (2024). Can You Tell the Difference? A Study of Human vs Machine-translated Subtitles. *Perspectives*, 32(6), 1115–1132. <https://doi.org/10.1080/0907676X.2023.2268149>
- Cerezo Merchán, B., Chaume Varela, F., Granell Zafra, J., Martí Ferriol, J. L., Martínez Sierra, J. J., & Marzà Ibàñez, A. (2016). *La traducción para el doblaje en España: mapa de convenciones*. Publicacions de la Universitat Jaume I. <https://doi.org/10.6035/estudistraduccio.trama.2016.3>
- Chaume, F. (2004). *Cine y traducción*. Cátedra.
- Chaume, F. (2007). Quality Standards in Dubbing: A Proposal. *TradTerm*, 13, 71–89. <https://doi.org/10.11606/issn.2317-9511.tradterm.2007.47466>
- Chaume, F. (2012). *Audiovisual Translation: Dubbing*. St. Jerome.
- Costa, C. B., & Silva, I. A. L. (2020). On the Translation of Literature as a Human Activity Par Excellence: Ethical Implications for Literary Machine Translation. *Aletria*, 30(4), 225–248. <https://doi.org/10.35699/2317-2096.2020.22047>
- de Higes Andino, I. (2023). Estudio del subtitulado automático bilingüe en la Comunitat Valenciana: el caso de L’Informatiu – Comunitat Valenciana de RTVE. In L. Mejías-Climent & J. de los

- Reyes Lozano (Eds.), *La traducción audiovisual a través de la traducción automática y la posesición: prácticas actuales y futuras* (pp. 95–112). Comares.
- de los Reyes Lozano, J., & Mejías-Climent, L. (2023a). La norma UNE-18587 sobre posesición y traducción automática: intersecciones con la industria del doblaje. In L. Mejías-Climent & J. de los Reyes Lozano (Eds.), *La traducción audiovisual a través de la traducción automática y la posesición: prácticas actuales y futuras* (pp. 81–93). Comares.
- de los Reyes Lozano, J., & Mejías-Climent, L. (2023b). Beyond the Black Mirror Effect: The Impact of Machine Translation in the Audiovisual Translation Environment. *Linguistica Antverpiensia*, 22, 1–19. <https://doi.org/10.52034/lans-tts.v22i.790>
- Díaz-Cintas, J. (2020). An Excursus on Audiovisual Translation. In L. Bogucki & M. Deckert (Eds.), *The Palgrave Handbook of Audiovisual Translation and Media Accessibility* (pp. 11–32). Palgrave Macmillan.
- Díaz Cintas, J., & Massidda, S. (2020). Technological Advances in Audiovisual Translation. In M. O'Hagan (Ed.), *The Routledge Handbook of Translation and Technology* (pp. 255–270). Routledge. <https://doi.org/10.4324/9781315311258-15>
- do Carmo, F. (2020). 'Time is Money' and the Value of Translation. *Translation Spaces*, 9(1), 35–57. <https://doi.org/10.1075/TS.00020.CAR>
- do Carmo, F., & Moorkens, J. (2022). Translation's New High-tech Clothes. In G. Massey, E. Huertas-Barros & D. Katan (Eds.), *The Human Translator in the 2020s* (pp. 11–26). Routledge. <https://doi.org/10.4324/9781003223344-2>
- Federico, M., Enyedi, R., Barra-Chicote, R., Giri, R., Isik, U., Krishnaswamy, A., & Sawaf, H. (2020). From Speech-to-Speech Translation to Automatic Dubbing. *arXiv preprint arXiv:2001.06785*. <https://doi.org/10.48550/arXiv.2001.06785>
- Flis, G., & Senda, T. (2022, September 15). What's in the Cards for AVT – Prognoses on the Future of Dubbing. *European Commission*. <https://knowledge-centre-translation-interpretation.ec.europa.eu/en/content/whats-cards-avt-prognoses-future-dubbing>
- Georgakopoulou, P. (2019). Technologization of Audiovisual Translation. In L. Pérez González (Ed.), *The Routledge Handbook of Audiovisual Translation* (pp. 516–539). Routledge.
- Georgakopoulou, P. (2020). The Faces of Audiovisual Translation. In European Parliament (Ed.), *The Many Faces of Translation: From Video Games to the Vatican* (pp. 16–33). European Parliament. <https://doi.org/10.2861/898773>
- Guerberof Arenas, A., & Toral, A. (2020). The Impact of Post-editing and Machine Translation on Creativity and Reading Experience. *Translation Spaces*, 9(2), 255–282. <https://doi.org/10.1075/ts.20035.gue>
- Görög, A. (2014). Quantifying and Benchmarking Quality: The TAUS Dynamic Quality Framework. *Tradumática*, (12), 443–454. <https://doi.org/10.5565/rev/tradumatica.66>
- Hadley, J. L., Taivalkoski-Shilov, K., Teixeira, C. S. C., & Toral, A. (Eds.). (2022). *Using Technologies for Creative-Text Translation*. Routledge.
- Hansen, D., & Houlmont, P.-Y. (2022). A Snapshot into the Possibility of Video Game Machine Translation. In J. Campbell, S. Larocca, J. Marciano, K. Savenkov & A. Yanishevsky (Eds.), *Proceedings of the 15th Conference of the Association for Machine Translation in the Americas* (Vol.

- 2; pp. 257–269). Association for Machine Translation in the Americas. <https://aclanthology.org/2022.amta-upg.18/>
- Hiraoka, Y., & Yamada, M. (2019). Pre-editing plus Neural Machine Translation for Subtitling: Effective Pre-editing Rules for Subtitling of TED Talks. In F. G. M. Forcada, A. Way, J. Tinsley, D. Shterionov, C. Rico (Eds.), *Proceedings of Machine Translation Summit XVII: Translator, Project and User Tracks* (pp. 64–72). European Association for Machine Translation. <https://aclanthology.org/W19-6710/>
- Karakanta, A., Bhattacharya, S., Nayak, S., Baumann, T., Negri, M., & Turchi, M. (2021). The Two Shades of Dubbing in Neural Machine Translation. In D. Scott, N. Bel, & C. Zong (Eds.), *Proceedings of the 28th International Conference on Computational Linguistics* (pp. 4327–4333). International Committee on Computational Linguistics. <https://doi.org/10.18653/v1/2020.coling-main.382>
- Koglin, A., Moura, W. H. C., Matos, M. A., & Silveira, J. G. P. (2023). Quality Assessment of Machine-translated Post-edited Subtitles: An Analysis of Brazilian Translators' Perceptions. *Linguistica Antverpiensia*, 22, 41–60. <https://doi.org/10.52034/lans-tts.v22i.765>
- Kress, G., & van Leeuwen, T. (2001). *Multimodal Discourse: The Modes and Media of Contemporary Communication*. Arnold.
- Lommel, A. (2018). Metrics for Translation Quality Assessment: A Case for Standardising Error Typologies. In J. Moorkens, S. Castilho, F. Gaspari, & S. Doherty (Eds.), *Translation Quality Assessment: From Principles to Practice* (pp. 109–127). Springer.
- Lommel, A., Uszkoreit, H., & Burchardt, A. (2014). Multidimensional Quality Metrics (MQM): A Framework for Declaring and Describing Translation Quality Metrics. *Tradumática*, (12), 455–463. <https://doi.org/10.5565/rev/tradumatica.77>
- Matousek, J., & Vít, J. (2012). Improving Automatic Dubbing with Subtitle Timing Optimisation Using Video Cut Detection. In *IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP), 2012* (pp. 2385–2388). IEEE. <https://doi.org/10.1109/ICASSP.2012.6288395>
- Mejías-Climent, L., & de los Reyes Lozano, J. (2021). Traducción automática y posesición en el aula de doblaje: resultados de una experiencia docente. *Hikma*, 20(2), 203–227. <https://doi.org/10.21071/hikma.v20i2.13383>
- Mejías-Climent, L., & Villanueva-Jordán, I. (2023). Traducción automática y audiovisual: la búsqueda del equilibrio entre optimización y procesamiento humano. In C. Rico Pérez & M. M. Sánchez Ramos (Eds.), *Traducción automática en contextos especializados* (pp. 209–231). Peter Lang.
- Miyata, R., & Fujita, A. (2021). Understanding Pre-editing for Black-box Neural Machine Translation. In P. Merlo, J. Tiedemann & R. Tsarfaty (Eds.), *Proceedings of the 16th Conference of the European Chapter of the Association for Computational Linguistics* (pp. 1539–1550). Association for Computational Linguistics. <https://doi.org/10.18653/v1/2021.eacl-main.132>
- Moniz, H., & Parra Escartín, C. (Eds.). (2023). *Towards Responsible Machine Translation: Ethical and Legal Considerations in Machine Translation*. Springer. <https://doi.org/10.1007/978-3-031-14689-3>
- O'Hagan, M. (Ed.). (2020). *The Routledge Handbook of Translation and Technology*. Routledge.
- Ortiz-Boix, C. (2016). Machine Translation and Post Editing in Wildlife Documentaries: Challenges and Possible Solutions. *Hermeneus*, 18, 269–313.

- Remael, A., & Revers, N. (2019). Multimodality and Audiovisual Translation: Cohesion in Accessible Films. In L. Pérez-González (Ed.), *The Routledge Handbook of Audiovisual Translation* (pp. 260–280). Routledge.
- Rico Pérez, C. (2023). Caminos convergentes en traducción audiovisual, traducción automática y posesición. In L. Mejías-Climent & J. de los Reyes Lozano (Eds.), *La traducción audiovisual a través de la traducción automática y la posesición: prácticas actuales y futuras* (pp. 1–14). Comares.
- Rico Pérez, C., & Sánchez Ramos, M. M. (2023). *Traducción automática en contextos especializados*. Peter Lang. <https://doi.org/10.3726/b20144>
- Rodríguez Fernández-Peña, A. C. (2023). Online Cloud Dubbing: How Home Recording Stormed the Dubbing Industry. *Tradumática*, (21), 28–48. <https://doi.org/10.5565/rev/tradumatica.335>
- Sakaguchi, N., Murawaki, Y., Chu, C., & Kurohashi, S. (2024). Identifying Source Language Expressions for Pre-editing in Machine Translation. In N. Calzolari, M.-Y. Kan, V. Hoste, A. Lenci, S. Sakti & N. Xue (Eds.), *Proceedings of the 2024 Joint International Conference on Computational Linguistics, Language Resources and Evaluation (LREC-COLING 2024)* (pp. 8605–8616). ELRA. <https://aclanthology.org/2024.lrec-main.755/>
- Sakamoto, A., van Laar, D., Moorkens, J., & do Carmo, F. (2024). Measuring Translators' Quality of Working Life and their Career Motivation: Conceptual and Methodological Aspects. *Translation Spaces*, 13(1), 54–77. <https://doi.org/10.1075/ts.23026.sak>
- Sánchez Ramos, M. M., & Rico Pérez, C. (2020). *Traducción Automática: conceptos clave, procesos de evaluación y técnicas de posesición*. Comares.
- Spiteri Miggiani, G. (2019). *Dialogue Writing for Dubbing: An Insider's Perspective*. Palgrave Macmillan.
- Spiteri Miggiani, G. (2022). Measuring Quality in Translation for Dubbing: A Quality Assessment Model Proposal for Trainers and Stakeholders. *XLinguae*, 15(2), 85–102. <https://doi.org/10.18355/XL.2022.15.02.07>
- Tam, D., Lakew, S. M., Virkar, Y., Mathur, P., & Federico, M. (2022). Isochrony-aware Neural Machine Translation for Automatic Dubbing. In *Proceedings of Interspeech 2022* (pp. 1776–1780). International Speech Communication Association. <https://doi.org/10.21437/Interspeech.2022-11136>
- Valdeón, R. A. (2022). Latest Trends in Audiovisual Translation. *Perspectives*, 30(3), 369–381. <https://doi.org/10.1080/0907676X.2022.2069226>
- Villanueva Jordán, I. A., & Romero-Muñoz, A. (2023). Nociones metodológicas para el análisis comparativo de traducciones automáticas para el doblaje. In L. Mejías-Climent & J. de los Reyes Lozano (Eds.), *La traducción audiovisual a través de la traducción automática y la posesición: prácticas actuales y futuras* (pp. 17–36). Comares.
- Vivziepop. (n.d.). *HELLUVA BOSS* [YouTube playlist]. YouTube. <https://www.youtube.com/playlist?list=PL-uopgYBi65HwiiDR9Y23lomAkGr9mm-S>
- Wang, L. (2023). Applying Automated Machine Translation to Educational Video Courses. *Education and Information Technologies*, 29(9), 10377–10390. <https://doi.org/10.1007/s10639-023-12219-0>

Yang, Y., Shillingford, B., Assael, Y., Wang, M., Liu, W., Chen, Y., Zhang, Y., Sezener, E., Cobo, L. C., Denil, M., Aytar, Y., & Freitas, N. D. (2020). Large Scale Multilingual Audio Visual Dubbing. *arXiv preprint arXiv:2011.03530*. <https://doi.org/10.48550/arXiv.2011.03530>

Zhang, X. (2025). A Review of Research on Pre-editing of Machine Translation. *Journal of Literature and Art Studies*, 15(1), 36–43. <https://doi.org/10.17265/2159-5836/2025.01.006>

Editorial notes

Authorship contribution

Conceptualization: L. Mejías-Climent

Data collection: A. Romero-Muñoz

Data analysis: A. Romero-Muñoz

Results and discussion: A. Romero-Muñoz, L. Mejías-Climent

Writing – review and editing: L. Mejías-Climent, A. Romero-Muñoz

Funding acquisition: L. Mejías-Climent

Research dataset

The research data is part of the funded DubTA project and was extracted from publicly available material on YouTube: Vivziepop. (n.d.). HELLUVA BOSS [YouTube playlist]. YouTube. <https://www.youtube.com/playlist?list=PL-uopgYBi65HwiiDR9Y23lomAkGr9mm-S>

Funding

DubTA. *La traducción automática aplicada a los procesos de traducción para el doblaje (Machine Translation Applied to Translation Processes for Dubbing)*. Funded by Universitat Jaume I (UJI-B2020-56). 2021-2022.

Image copyright

Not applicable.

Approval by ethics committee

Not applicable.

Conflicts of interest

Not applicable.

Data availability statement

The data from this research, which are not included in this work, may be made available by the author(s) upon request.

License

The authors grant *Cadernos de Tradução* exclusive rights for first publication, while simultaneously licensing the work under the [Creative Commons Attribution 4.0 International \(CC BY 4.0\)](https://creativecommons.org/licenses/by/4.0/) License. This license enables third parties to remix, adapt, and create from the published work, while giving proper credit to the authors and acknowledging the initial publication in this journal. Authors are permitted to enter into additional agreements separately for the non-exclusive distribution of the published version of the work in this journal. This may include publishing it in an institutional repository, on a personal website, on academic social networks, publishing a translation, or republishing the work as a book chapter, all with due recognition of authorship and first publication in this journal.

Publisher

Cadernos de Tradução is a publication of the Graduate Program in Translation Studies at the Federal University of Santa Catarina. The journal *Cadernos de Tradução* is hosted by the [Portal de Periódicos UFSC](https://portal.periodicos.ufsc.br/). The ideas expressed in this paper are the responsibility of its authors and do not necessarily represent the views of the editors or the university.

Section editors

Andréia Guerini – Willian Moura



Cadernos de Tradução, 46, 2026, e105649
Graduate Program in Translation Studies
Federal University of Santa Catarina, Brazil. ISSN 2175-7968
DOI <https://doi.org/10.5007/2175-7968.2026.e105649>

Style editors

Alice S. Rezende – Ingrid Bignardi – João G. P. Silveira – Kamila Oliveira

Article history

Received: 11-08-2025

Approved: 26-01-2026

Revised: 20-02-2026

Published: 03-2026



Cadernos de Tradução, 46, 2026, e105649

Graduate Program in Translation Studies

Federal University of Santa Catarina, Brazil. ISSN 2175-7968

DOI <https://doi.org/10.5007/2175-7968.2026.e105649>