



**UNIVERSIDADE FEDERAL DO CEARÁ**  
**CENTRO DE HUMANIDADES**  
**PRÓ-REITORIA DE PESQUISA E PÓS-GRADUAÇÃO**  
**PROGRAMA DE PÓS-GRADUAÇÃO EM ESTUDOS DA TRADUÇÃO**

**ALBERTO HOLANDA PIMENTEL NETO**

**MOVIMENTAÇÃO OCULAR DURANTE A LEITURA DE**  
**LEGENDAS INTRALINGUÍSTICAS EM VÍDEOS COM TEXTO EM**  
**TELA: ESTUDO PILOTO EXPERIMENTAL**

**FORTALEZA**

**2024**

ALBERTO HOLANDA PIMENTEL NETO

MOVIMENTAÇÃO OCULAR DURANTE A LEITURA DE  
LEGENDAS INTRALINGUÍSTICAS EM VÍDEOS COM TEXTO EM  
TELA: ESTUDO PILOTO EXPERIMENTAL

Dissertação apresentada ao Programa de Pós- Graduação em Estudos da Tradução (POET) da Universidade Federal do Ceará, como requisito parcial à obtenção do título de Mestre. Área de concentração: Estudos da Tradução.

Orientadora: Prof.<sup>a</sup> Dra. Maria Cristina Micelli Fonseca.

**FORTALEZA**

**2024**

Dados Internacionais de Catalogação na Publicação  
Universidade Federal do Ceará  
Sistema de Bibliotecas  
Gerada automaticamente pelo módulo Catalog, mediante os dados fornecidos pelo(a) autor(a)

---

P698m      Pimentel Neto, Alberto Holanda.  
Movimentação ocular durante a leitura de legendas intralinguísticas em vídeos com texto em tela : estudo piloto experimental / Alberto Holanda Pimentel Neto. – 2024.  
181 f. : il. color.

Dissertação (mestrado) – Universidade Federal do Ceará, Centro de Humanidades, Programa de PósGraduação em Letras, Fortaleza, 2024.  
Orientação: Profa. Dra. Maria Cristina Micelli Fonseca.

1. Estudos da tradução. 2. Legendagem intralinguística. 3. Movimentação ocular. 4. Vídeos legendados. 5. Texto em tela. I. Título.

CDD 400

---

ALBERTO HOLANDA PIMENTEL NETO

MOVIMENTAÇÃO OCULAR DURANTE A LEITURA DE  
LEGENDAS INTRALINGUÍSTICAS EM VÍDEOS COM TEXTO EM  
TELA: ESTUDO PILOTO EXPERIMENTAL

Dissertação apresentada ao Programa de Pós-Graduação em Estudos da Tradução (POET) da Universidade Federal do Ceará, como requisito parcial à obtenção do título de Mestre. Área de concentração: Estudos da Tradução.

Aprovada em: 30 / 08 / 2024.

BANCA EXAMINADORA

---

Profa. Dra. Maria Cristina Micelli Fonseca (Orientadora)  
Universidade Federal do Ceará (UFC)

---

Profa. Dra. Elisângela Nogueira Teixeira  
Universidade Federal do Ceará (UFC)

---

Prof. Dr. Ítalo Alves Pinto de Assis  
Universidade Estadual Vale do Acaraú (UVA)

---

Profa. Dra. Erica dos Santos Rodrigues  
Pontifícia Universidade Católica do Rio de Janeiro (PUC-Rio)

À minha mãe, minha base.

## **AGRADECIMENTOS**

Em geral, agradecimentos são sucintos e diretos. Bom, meu percurso ao longo do Mestrado não foi. Foram vários altos e baixos, vários meses de trajetória, várias idas e voltas e, é claro, várias pessoas que fizeram parte de tudo isso e acompanharam minha luta para chegar até aqui (sim, luta, pois não foi fácil). Então, seria impossível ser sucinto em meus agradecimentos. Portanto, segue.

PS. Não espere um estilo formal nem tradicional aqui. Afinal de contas, a própria pesquisa vai de encontro a isso. ;)

A Deus, antes de tudo.

A mim, por ter persistido até aqui, é claro, né?

À minha mãe, pelo apoio, compreensão, qualidade humana, ética e espelho de vida. Jamais terei como agradecer o suficiente.

À minha orientadora, Profa. Maria Cristina, pela paciência, por aceitar me orientar e abraçar o desafio de recomeçar toda uma pesquisa do zero no meio do caminho.

À FUNCAP, pelo financiamento.

Ao Laboratório de Aquisição, Desenvolvimento e Processamento da Linguagem, por disponibilizar o espaço para a coleta de dados.

A Priscila, por me motivar a iniciar o Mestrado.

Aos participantes que aceitaram participar da pesquisa. Aos meus colegas de curso na POET.

Aos professores da POET e a todos os que passaram por minha trajetória acadêmica até aqui.

À Profa. Camila Braga, por ceder suas aulas para meu estágio de docência e testemunhar meu processo.

À banca examinadora, por aceitar o convite e contribuir com as discussões do trabalho.

A Karol e Dherek, pelos “dissertados” e por sempre apoiarem. Nossos grupos de estudo foram essenciais.

A Israel e Filipe, pela ajuda com os rascunhos iniciais nas legendas.

Um agradecimento especial a Filipe, pela ajuda em diversos seguimentos, além das legendas, pela ajuda com a seleção de vídeos, pelo apoio e

suporte pessoal e, imensamente, pelo suporte pessoal e profissional. Isso foi essencial na reta final da escrita.

A todos os meus colegas de profissão, pelo suporte. A Pilar Camacho, por toda a torcida e motivação.

A Ian, Angela, Dallyana, Andreina, e em especial, a Lueuda, por compartilharem das dificuldades da Pós dentro e fora da POET.

A Karol, novamente, e Arthur, por estarem abertos a ouvir e estarem sempre por perto. A Thiago Dantas, pelas ajudas de última hora.

A meu velho amigo, Jayme Welton, com quem posso sempre contar e por se disponibilizar a ficar com minha mãe quando precisei viajar à Fortaleza.

A Samara Myria e Sara Karlessandra, por me receberem e me darem o apoio de que eu precisei em Fortaleza. Um agradecimento especial a “Pessoa” (Samara), como carinhosamente nos tratamos, por aceitar ser a atriz dos vídeos do experimento.

A Iana, por todo o apoio inicial no Laboratório e apoio pessoal ao longo de toda a montagem do experimento e coleta de dados. Seu “vai dar certo” e “vai dar bom” realmente deu certo. Esperemos que dê bom. Muito obrigado!

A Vinícius, pelo imenso apoio com o *Experiment Builder*.

A Manu, por me ajudar com a coleta e os acessos ao *Data Viewer*. A João, pelos *insights* sobre os dados.

Um IMENSO AGRADECIMENTO (em letras maiúsculas, sim) ao suporte do *SR Research Support*, mais especificamente a Marcus Johnson, por ter sido o ponto de contato com o suporte e quem me acompanhou e orientou passo a passo desde o início com todos os aspectos e dificuldades enfrentadas com *Eye Link*. Sem você, nada disso teria sido possível. Meu eterno obrigado!

Ao meu psicólogo e psiquiatra, por todas as sessões de terapia. Sim, foi necessário e indispensável. Tenho que agradecer, né? :)

Por fim, um pedido de desculpas caso eu tenha deixado de agradecer a alguém! A todos vocês, meu muito obrigado!

“Que as tuas palavras sejam  
minhas legendas neste mundo  
estrangeiro a mim”.

Próprio autor

## RESUMO

Este trabalho objetiva-se a investigar como o design de legendas em vídeos com texto em tela influencia o processamento de leitura. A legendagem, no que tange aos Estudos da Tradução, é um tipo de tradução audiovisual. Dependendo do tipo de legenda produzida, esta pode ser considerada como uma forma de tradução intralinguística, interlinguística ou intersemiótica. É comum que essas categorias de legendas apresentem narrativas de natureza imagética, conhecidas dentro do mercado de legendagem como "*on-screen texts*", ou textos em tela, cujo uso é bastante comum e, muitas vezes, competem com a informação de legenda que é apresentada. Para isso, preparou-se um experimento com 6 vídeos em duas condições experimentais: condição 1 (legendas totais) e condição 2 (legendas parciais). O experimento foi construído no *Experiment Builder* com o rastreador ocular *EyeLink 1000 Plus*. Para a exibição dos vídeos, foram montadas 4 listas aleatórias (A, B, C e D), para garantir que a ordem de exibição dos vídeos não influenciasse os resultados, e os dados foram analisados no *Data Viewer*. A análise foi realizada através de métodos de rastreamento ocular, focando em métricas baseadas na fixação dos participantes em ambas as áreas de interesse, e resposta ao questionário pós-coleta, com perguntas sobre o conteúdo dos vídeos. Os resultados inclinam-se na direção de que legendas parciais podem ser mais eficazes em contextos onde há texto em tela redundante, sugerindo que um design de legendas que evite redundâncias e otimize a apresentação da informação visual pode melhorar a experiência de visualização e a compreensão do espectador. Espera-se que os resultados obtidos possam fornecer dados válidos para a construção metodológica de pesquisas com foco na movimentação ocular durante a leitura de legendas em vídeos com textos em tela.

**Palavras-chave:** estudos da tradução; legendagem intralinguística; movimentação ocular; vídeos legendados; texto em tela.

## **ABSTRACT**

This thesis aims at investigating how the design of subtitles in videos with on-screen texts influences reading processing. Subtitling, within Translation Studies, is a type of audiovisual translation. Depending on the type of subtitle produced, it can be considered a form of intralinguistic, interlinguistic, or intersemiotic translation. These categories of subtitles commonly include image-based narratives, known within the subtitling industry as "on-screen texts.", whose use is quite common, and they often compete with the subtitle information presented. For this purpose, an experiment was prepared with 6 videos in two experimental conditions: condition 1 (total subtitles) and condition 2 (partial subtitles). The experiment was built in Experiment Builder using the EyeLink 1000 Plus eye tracker. Four random lists (A, B, C, and D) were assembled to display the videos, ensuring that the order of video display does not influence the results, and the data were analyzed in Data Viewer. The analysis was conducted through eye-tracking methods, focusing on metrics based on the participants' fixation on both areas of interest and responses to a post-collection questionnaire about the video content. The results lean towards the direction that partial subtitles may be more effective in contexts where there is redundant on-screen text, suggesting that a subtitle design that avoids redundancies and optimizes the presentation of visual information can enhance the viewing experience and viewer comprehension. It is hoped that the results obtained can provide valid data for the methodological construction of research focused on eye movement during the reading of subtitles in videos with on-screen texts.

**Keywords:** translation studies; intralingual subtitling; eye tracking; captioned Videos; on-screen texts.

## LISTA DE FIGURAS

|             |                                                                                                        |    |
|-------------|--------------------------------------------------------------------------------------------------------|----|
| Figura 1 -  | Esquema ilustrativo de produção de legendas a partir de duas fontes.....                               | 33 |
| Figura 2 -  | Texto em tela: ID Gráfico de localização traduzido com Forced Narratives.....                          | 37 |
| Figura 3 -  | Estratégia de Localização em “Divertidamente” por Edição de Vídeo.....                                 | 40 |
| Figura 4 -  | Estratégia de Localização 2 em “Divertidamente” por edição de vídeo.....                               | 40 |
| Figura 5 -  | Captura de tela: tradução por edição de vídeo para o texto em tela em vídeo do canal FortNine.....     | 41 |
| Figura 6 -  | Captura de tela: Tradução por edição de vídeo para o texto em tela em vídeo do canal Felipe Neto ..... | 41 |
| Figura 7 -  | Captura de tela: texto em tela através de Forced Narratives em “O Impossível” (1).....                 | 42 |
| Figura 8 -  | Captura de tela: texto em tela através de Forced Narratives em “O Impossível” (2).....                 | 43 |
| Figura 9 -  | Captura de tela: texto em tela com narração em “Harry Potter e a Câmara Secreta”.....                  | 44 |
| Figura 10 - | Captura de tela: Texto em Tela com narração em material intralinguístico.....                          | 45 |
| Figura 11 - | Captura de tela: Exemplo de Legenda Total .....                                                        | 47 |
| Figura 12 - | Captura de Tela: Sequência de imagens como exemplo de Legenda Parcial 1.....                           | 49 |
| Figura 13 - | Captura de Tela: Sequência de imagens como exemplo de Legenda Parcial 2.....                           | 49 |
| Figura 14 - | Captura de Tela: Sequência de imagens como exemplo de Legenda Parcial 3.....                           | 50 |
| Figura 15 - | Captura de Tela: Sequência de imagens como exemplo de Legenda Parcial e Total.....                     | 51 |
| Figura 16 - | Captura de Tela: Sequência de imagens como exemplo de                                                  |    |

|             |                                                                                                                                                         |     |
|-------------|---------------------------------------------------------------------------------------------------------------------------------------------------------|-----|
|             | legenda em texto em tela dinâmico.....                                                                                                                  | 52  |
| Figura 17 - | Condições experimentais para cada vídeo da pesquisa de Liao, Kruger e Doherty (2020).....                                                               | 55  |
| Figura 18 - | Áreas de Interesse definidas por Liao, Kruger e Doherty (2020), em três regiões do vídeo.....                                                           | 56  |
| Figura 19 - | Exemplo de estímulo usado por Liao et al. (2021), com a presença (imagem à direita) e a ausência (imagem à esquerda) de informação visual de fundo..... | 57  |
| Figura 20 - | Áreas de Interesse definidas por O'Hagan e Sasamoto (2016) .....                                                                                        | 60  |
| Figura 21 - | Distribuição de duração de fixações .....                                                                                                               | 63  |
| Figura 22 - | Exemplo de sequência de fixações em função das áreas de interesse (traduzido).....                                                                      | 66  |
| Figura 23 - | Representação imagética dos textos em tela utilizados.....                                                                                              | 85  |
| Figura 24 - | Sequência de imagens para período crítico no vídeo 1 na condição total.....                                                                             | 87  |
| Figura 25 - | Captura de Tela: Sequência de imagens para período crítico no vídeo 1 na condição parcial.....                                                          | 87  |
| Figura 26 - | Interface do TRIAL_Datasource no Experiment Builder.....                                                                                                | 92  |
| Figura 27 - | Criação de fórmula para códigos dos participantes.....                                                                                                  | 95  |
| Figura 28 - | Áreas de interesse delimitadas para os vídeos.....                                                                                                      | 99  |
| Figura 29 - | Período de Interesse Principal, “fullvideos”, no “Interest Period Editor” .....                                                                         | 100 |
| Figura 30 - | Lista de eventos antes e após adição de mensagens.....                                                                                                  | 101 |
| Figura 31 - | Exemplo de código de mensagem na lista de eventos.....                                                                                                  | 102 |
| Figura 32 - | Exemplo de comando de importação de mensagens pré-definidas.....                                                                                        | 105 |
| Figura 33 - | Criação de período de interesse no editor.....                                                                                                          | 106 |

|             |                                                                                                         |     |
|-------------|---------------------------------------------------------------------------------------------------------|-----|
| Figura 34 - | Exemplo de mapa de fixações antes e depois.....                                                         | 107 |
| Figura 35 - | Sequência de mapas de calor para a condição 1, “Legendas Totais”, para cada período de interesse.....   | 112 |
| Figura 36 - | Sequência de mapas de calor para a condição 2, “Legendas Parciais”, para cada período de interesse..... | 112 |
| Figura 37 - | Sistema de pontuação por sujeito para o questionário pós-coleta .....                                   | 128 |
| Figura 38 - | Sequência de Fixações de LDP16.....                                                                     | 143 |
| Figura 39 - | Sequência de Fixações de LDP20 para CP3 no vídeo 4 na condição 1 .....                                  | 144 |
| Figura 40 - | Sequência de Fixações de LDP20 para CP1 no vídeo 4 na condição 1.....                                   | 145 |

## LISTA DE GRÁFICOS

|              |                                                                                                    |     |  |
|--------------|----------------------------------------------------------------------------------------------------|-----|--|
| Gráfico 1 -  | Tempo Médio de Permanência em cada Área de Interesse por Participante em ambas as condições.....   | 110 |  |
| Gráfico 2 -  | Representação estatística do Tempo Médio de Permanência por Condição Experimental.....             | 113 |  |
| Gráfico 3 -  | Representação estatística do Tempo Médio de Permanência por Área de Interesse.....                 | 113 |  |
| Gráfico 4 -  | Duração Média de Fixação por Participante em Ambas as Condições.....                               | 114 |  |
| Gráfico 5 -  | Duração Média das Fixações por participante na condição 1.....                                     | 115 |  |
| Gráfico 6 -  | Duração Média das Fixações por participante na condição.....                                       | 116 |  |
| Gráfico 7 -  | Dados Estatísticos para a Duração Média de Fixações por Condição Experimental.....                 | 117 |  |
| Gráfico 8 -  | Dados Estatísticos para a Duração Média de Fixações por Al...                                      | 117 |  |
| Gráfico 9 -  | Dados Estatísticos para a Média de Fixações por Condição Experimental.....                         | 118 |  |
| Gráfico 10 - | Média de Fixações por Área de Interesse.....                                                       | 119 |  |
| Gráfico 11 - | Duração média da primeira fixação por participante em ambas as condições.....                      | 120 |  |
| Gráfico 12 - | Duração média da primeira fixação por participante na condição 1.....                              | 121 |  |
| Gráfico 13 - | Duração Média da Primeira Fixação por participante na condição 2.....                              | 121 |  |
| Gráfico 14 - | Média da Primeira fixação por período de interesse na condição 1.....                              | 122 |  |
| Gráfico 15 - | Média da Primeira fixação por período de interesse na condição 2.....                              | 122 |  |
| Gráfico 16 - | Dados estatísticos para a duração média da primeira fixação nas Als por condição experimental..... | 123 |  |
| Gráfico 17 - | Dados estatísticos para a duração média da primeira fixação nas Als por Área de Interesse.....     | 124 |  |

## LISTA DE QUADROS

|             |                                                                                                                |     |  |
|-------------|----------------------------------------------------------------------------------------------------------------|-----|--|
| Quadro 1 -  | Conversão de valores entre palavras por minuto (ppm) e caracteres por segundo incluindo espaços em branco..... | 32  |  |
| Quadro 2 -  | Conversão de valores entre palavras por minuto (ppm) e caracteres por segundo excluindo espaços em branco..... | 32  |  |
| Quadro 3 -  | Duração média das fixações em leitura, busca visual e percepção de cena (adaptado).....                        | 64  |  |
| Quadro 4 -  | Matriz de fixações entre as áreas de interesse da Fig 22.....                                                  | 66  |  |
| Quadro 5 -  | Matriz de fixações entre AI 6 e AI 7 adaptado de Eyelink (2024).....                                           | 67  |  |
| Quadro 6 -  | Sequência de legendas para período crítico do vídeo 1.....                                                     | 86  |  |
| Quadro 7 -  | Perfil Socioeconômico dos participantes em idade, sexo e nível de escolaridade.....                            | 96  |  |
| Quadro 8 -  | Perfil dos participantes em relação aos hábitos de acesso a redes sociais.....                                 | 97  |  |
| Quadro 9 -  | Sequência de vídeos e suas condições experimentais por lista assistida.....                                    | 97  |  |
| Quadro 10 - | Exemplo de lista de mensagem para a lista de visualização D....                                                | 103 |  |
| Quadro 11 - | Afirmativas erradas por participante e vídeo relacionado.....                                                  | 129 |  |

## LISTA DE TABELAS

|             |                                                               |    |  |
|-------------|---------------------------------------------------------------|----|--|
| Tabela 1 -  | Número de palavras dos roteiros .....                         | 72 |  |
| Tabela 2 -  | Frequência das principais palavras de cada roteiro .....      | 73 |  |
| Tabela 3 -  | Média da frequência e tamanho das palavras dos roteiros ..... | 74 |  |
| Tabela 4 -  | Número de palavras dos roteiros .....                         | 75 |  |
| Tabela 5 -  | Valores de CPS médio para os vídeos preparados .....          | 75 |  |
| Tabela 6 -  | Valores médios dos fatores de influência.....                 | 76 |  |
| Tabela 7 -  | Texto em tela 1 para o vídeo 1 .....                          | 77 |  |
| Tabela 8 -  | Texto em tela 2 para o vídeo 1 .....                          | 77 |  |
| Tabela 9 -  | Texto em tela 3 para o vídeo 1 .....                          | 77 |  |
| Tabela 10 - | Texto em tela 4 para o vídeo 1 .....                          | 78 |  |
| Tabela 11 - | Texto em tela 1 para o vídeo 2 .....                          | 78 |  |
| Tabela 12 - | Texto em tela 2 para o vídeo 2 .....                          | 78 |  |
| Tabela 13 - | Texto em tela 3 para o vídeo 2 .....                          | 79 |  |
| Tabela 14 - | Texto em tela 4 para o vídeo 2 .....                          | 79 |  |
| Tabela 15 - | Texto em tela 1 para o vídeo 3 .....                          | 79 |  |
| Tabela 16 - | Texto em tela 2 para o vídeo 3 .....                          | 80 |  |
| Tabela 17 - | Texto em tela 3 para o vídeo 3 .....                          | 80 |  |
| Tabela 18 - | Texto em tela 4 para o vídeo 3.....                           | 80 |  |
| Tabela 19 - | Texto em tela 1 para o vídeo 4.....                           | 81 |  |
| Tabela 20 - | Texto em tela 2 para o vídeo 4.....                           | 81 |  |
| Tabela 21 - | Texto em tela 3 para o vídeo 4.....                           | 81 |  |
| Tabela 22 - | Texto em tela 4 para o vídeo 4.....                           | 82 |  |
| Tabela 23 - | Texto em tela 1 para o vídeo 5.....                           | 82 |  |
| Tabela 24 - | Texto em tela 2 para o vídeo 5.....                           | 82 |  |
| Tabela 25 - | Texto em tela 3 para o vídeo 5.....                           | 83 |  |
| Tabela 26 - | Texto em tela 4 para o vídeo 5.....                           | 83 |  |
| Tabela 27 - | Texto em tela 1 para o vídeo 6.....                           | 83 |  |
| Tabela 28 - | Texto em tela 2 para o vídeo 6.....                           | 84 |  |
| Tabela 29 - | Texto em tela 3 para o vídeo 6.....                           | 84 |  |
| Tabela 30 - | Texto em tela 4 para o vídeo 6.....                           | 84 |  |
| Tabela 31 - | Lista aleatória A .....                                       | 89 |  |

|             |                                                                                                            |     |  |
|-------------|------------------------------------------------------------------------------------------------------------|-----|--|
| Tabela 32 - | Tempo médio de permanência em cada Área de Interesse em ambas as condições experimentais .....             | 110 |  |
| Tabela 33 - | Tempo Médio de Permanência nas AI por Condição Experimental .....                                          | 111 |  |
| Tabela 34 - | Duração média de fixações para as Áreas de Interesse AI 1 e AI 2 em ambas as condições experimentais ..... | 114 |  |
| Tabela 35 - | Duração Média de Fixações por AI e por condição experimental .....                                         | 116 |  |
| Tabela 36 - | Média de fixações por AI e por condição experimental .....                                                 | 118 |  |
| Tabela 37 - | Duração Média da Primeira Fixação por IA em ambas as condições .....                                       | 120 |  |
| Tabela 38 - | Duração média da primeira fixação nas AIs por condição experimental .....                                  | 123 |  |
| Tabela 39 - | Quantidade média de fixações entre áreas de interesse por participante na condição 1 .....                 | 125 |  |
| Tabela 40 - | Quantidade média de fixações entre áreas de interesse por participante, na condição 2 .....                | 126 |  |
| Tabela 41 - | Quantidade média de fixações entre áreas de interesse por área de interesse na condição 1 .....            | 127 |  |
| Tabela 42 - | Quantidade média de fixações entre áreas de interesse por área de interesse na condição 2 .....            | 127 |  |
| Tabela 43 - | Total de acertos por participante .....                                                                    | 128 |  |
| Tabela 44 - | Medidas de Fixação de LDP16 na condição 1 .....                                                            | 135 |  |
| Tabela 45 - | Valores médio de duração, número de caracteres e caracteres/segundo por período de interesse .....         | 138 |  |

## SUMÁRIO

|              |                                                                                                       |           |
|--------------|-------------------------------------------------------------------------------------------------------|-----------|
| <b>1</b>     | <b>INTRODUÇÃO .....</b>                                                                               | <b>20</b> |
| <b>2</b>     | <b>FUNDAMENTAÇÃO TEÓRICA .....</b>                                                                    | <b>26</b> |
| <b>2.1</b>   | <b>A legendagem tradicional na tradução audiovisual .....</b>                                         | <b>26</b> |
| <b>2.1.1</b> | <b><i>Classificação .....</i></b>                                                                     | <b>27</b> |
| 2.1.1.1      | <i>Tempo disponível para preparação .....</i>                                                         | 29        |
| 2.1.1.2      | <i>Modos de exibição .....</i>                                                                        | 29        |
| 2.1.1.3      | <i>Métodos de projeção e meios de distribuição .....</i>                                              | 30        |
| <b>2.1.2</b> | <b><i>A legendagem para o público em geral .....</i></b>                                              | <b>31</b> |
| <b>2.1.3</b> | <b><i>Por que o uso de legendas intralinguísticas? .....</i></b>                                      | <b>34</b> |
| <b>2.2</b>   | <b>On-Screen Texts (Textos em tela) .....</b>                                                         | <b>35</b> |
| <b>2.2.1</b> | <b><i>Exemplos de textos em tela interlinguísticos através de<br/>Localização Imagética .....</i></b> | <b>39</b> |
| <b>2.2.2</b> | <b><i>Exemplos de textos em tela interlinguísticos através de Forced<br/>Narrativas .....</i></b>     | <b>42</b> |
| <b>2.2.3</b> | <b><i>Exemplos de textos em tela com narração intralinguística .....</i></b>                          | <b>44</b> |
| <b>2.2.4</b> | <b><i>O uso de legendas e textos em telas nas redes sociais .....</i></b>                             | <b>46</b> |
| 2.2.4.1      | <i>Exemplos de legendas totais .....</i>                                                              | 47        |
| 2.2.4.2      | <i>Exemplos de legendas parciais .....</i>                                                            | 48        |
| 2.2.4.3      | <i>Outros exemplos .....</i>                                                                          | 50        |
| <b>2.3</b>   | <b>Rastreamento ocular e legendagem .....</b>                                                         | <b>53</b> |
| <b>2.4</b>   | <b>Hipóteses .....</b>                                                                                | <b>62</b> |
| <b>3</b>     | <b>DESENHO EXPERIMENTAL .....</b>                                                                     | <b>70</b> |
| <b>3.1</b>   | <b>Tipo de pesquisa .....</b>                                                                         | <b>71</b> |
| <b>3.2</b>   | <b>Preparação do Material .....</b>                                                                   | <b>71</b> |
| <b>3.2.1</b> | <b><i>Seleção dos vídeos .....</i></b>                                                                | <b>71</b> |
| <b>3.2.2</b> | <b><i>Controle de variáveis nos vídeos .....</i></b>                                                  | <b>72</b> |
| <b>3.2.3</b> | <b><i>Gravação dos vídeos .....</i></b>                                                               | <b>74</b> |
| <b>3.2.4</b> | <b><i>Preparação das legendas .....</i></b>                                                           | <b>75</b> |
| <b>3.2.5</b> | <b><i>Delimitação dos textos em tela .....</i></b>                                                    | <b>76</b> |
| 3.2.5.1      | <i>Legendas manipuladas para o vídeo 1 .....</i>                                                      | 77        |
| 3.2.5.2      | <i>Legendas manipuladas para o vídeo 2 .....</i>                                                      | 78        |
| 3.2.5.3      | <i>Legendas manipuladas para o vídeo 3 .....</i>                                                      | 79        |

|              |                                                               |            |
|--------------|---------------------------------------------------------------|------------|
| 3.2.5.4      | <i>Legendas manipuladas para o vídeo 4 ..</i>                 | 81         |
| 3.2.5.5      | <i>Legendas manipuladas para o vídeo 5 ..</i>                 | 82         |
| 3.2.5.6      | <i>Legendas manipuladas para o vídeo 6 ..</i>                 | 83         |
| <b>3.2.6</b> | <b>Renderização dos textos em tela ..</b>                     | <b>85</b>  |
| 3.2.6.1      | <i>Textos em tela ..</i>                                      | 85         |
| <b>3.3</b>   | <b>Condições experimentais ..</b>                             | <b>86</b>  |
| <b>3.3.1</b> | <b><i>Condição 1: legendas totais ..</i></b>                  | <b>86</b>  |
| <b>3.3.2</b> | <b><i>Condição 2: legendas parciais ..</i></b>                | <b>87</b>  |
| <b>3.4</b>   | <b>Experiment Builder ..</b>                                  | <b>88</b>  |
| <b>3.5</b>   | <b>Áreas de interesse ..</b>                                  | <b>92</b>  |
| <b>3.6</b>   | <b>Aspectos éticos da pesquisa ..</b>                         | <b>93</b>  |
| <b>3.7</b>   | <b>Etapas de coleta ..</b>                                    | <b>94</b>  |
| <b>3.7.1</b> | <b><i>Procedimentos pré-coleta ..</i></b>                     | <b>94</b>  |
| <b>3.7.2</b> | <b><i>Procedimentos de coleta ..</i></b>                      | <b>94</b>  |
| <b>3.7.3</b> | <b><i>Participantes para a coleta ..</i></b>                  | <b>95</b>  |
| 3.7.3.1      | <i>Listas assistidas ..</i>                                   | 97         |
| <b>3.8</b>   | <b>Tratamento dos dados da movimentação ocular ..</b>         | <b>98</b>  |
| <b>3.8.1</b> | <b><i>Áreas de interesse ..</i></b>                           | <b>100</b> |
| 3.8.1.2      | <i>Períodos de interesse ..</i>                               | 109        |
| <b>3.8.2</b> | <b><i>Medidas de análise ..</i></b>                           | <b>107</b> |
| <b>4</b>     | <b>RESULTADOS ..</b>                                          | <b>109</b> |
| <b>4.1</b>   | <b>Tempo de permanência na área de interesse ..</b>           | <b>109</b> |
| <b>4.2</b>   | <b>Médias de fixações por área de interesse ..</b>            | <b>114</b> |
| <b>4.3</b>   | <b>Medidas de fixações por condição experimental ..</b>       | <b>115</b> |
| <b>4.4</b>   | <b>Tempo da primeira fixação na área de interesse ..</b>      | <b>119</b> |
| <b>4.5</b>   | <b>Análise de sequência das fixações ..</b>                   | <b>124</b> |
| <b>4.6</b>   | <b>Questionário pós-coleta ..</b>                             | <b>128</b> |
| <b>5</b>     | <b>DISCUSSÕES INICIAIS ..</b>                                 | <b>130</b> |
| <b>5.1</b>   | <b>Tempo de permanência na área de interesse ..</b>           | <b>130</b> |
| <b>5.2</b>   | <b>Medidas médias de fixações ..</b>                          | <b>132</b> |
| <b>5.2.1</b> | <b><i>Médias de fixações por condição experimental ..</i></b> | <b>134</b> |
| <b>5.2.2</b> | <b><i>Médias de fixação por período de interesse ..</i></b>   | <b>136</b> |

|     |                                                                  |     |
|-----|------------------------------------------------------------------|-----|
| 5.3 | Tempo da primeira fixação na área de interesse .....             | 139 |
| 5.4 | Análise de sequência das fixações .....                          | 141 |
| 5.5 | Questionário pós-coleta .....                                    | 143 |
| 6   | CONSIDERAÇÕES FINAIS .....                                       | 147 |
|     | REFERÊNCIAS .....                                                | 153 |
|     | APÊNDICE A – Questionário pré-coleta .....                       | 160 |
|     | APÊNDICE B – Questionário pós-coleta .....                       | 165 |
|     | APÊNDICE C – Roteiro para o vídeo 1 .....                        | 167 |
|     | APÊNDICE D – Roteiro para o vídeo 2 .....                        | 168 |
|     | APÊNDICE E – Roteiro para o vídeo 3 .....                        | 169 |
|     | APÊNDICE F – Roteiro para o vídeo 4 .....                        | 170 |
|     | APÊNDICE G – Roteiro para o vídeo 5 .....                        | 171 |
|     | APÊNDICE H – Roteiro para o vídeo 6 .....                        | 172 |
|     | APÊNDICE I – Termo de Consentimento de Uso de Imagem e Voz ..... | 173 |
|     | APÊNDICE J – Termo de Consentimento Livre e Esclarecido .....    | 174 |
|     | APÊNDICE K – Mensagens da lista A .....                          | 176 |
|     | APÊNDICE L – Mensagens da lista D .....                          | 178 |
|     | ANEXO A – Parecer Comitê de Ética .....                          | 180 |

## 1 INTRODUÇÃO

Impessoalidade. É o que dita e preza-se na escrita acadêmica. No entanto, seria impossível, se não hipocrisia minha, escrever a introdução deste trabalho sem me colocar nela enquanto sujeito atuante em cada passo que a compôs. A começar pelo próprio tema, cujos motivos se originaram a partir de minha experiência enquanto tradutor, em particular, como legendista. Dessa forma, julgo necessário descrever, brevemente, minha experiência profissional para que se compreenda, então, os objetivos acadêmicos que me fizeram iniciar esta pesquisa.

Sou tradutor legendista há seis anos, no momento de escrita desta dissertação. E, durante esse breve período de trabalho, tive a oportunidade de atuar profissionalmente em várias escalas da produção de conteúdo audiovisual, desde a posição de legendista freelancer a gerente de projetos audiovisuais, e com um número razoável de clientes. Por motivos éticos, não tenho permissão para expor ou detalhar normas, recomendações ou mesmo diretrizes específicas de clientes/agências. No entanto, dentre as diferenças entre plataformas, produtoras e agências de tradução voltadas ao mercado audiovisual, sempre me intrigou a baixa, ou (arrisco dizer) inexistente, padronização no que tange às recomendações técnicas para vídeos voltados às redes sociais ou à divulgação em plataformas abertas, como YouTube.

Sendo assim, por ora, trago apenas os fatores que fundamentam a escolha do meu tema de pesquisa, cujo interesse inicial de investigação surgiu a partir da minha própria experiência de trabalho no contexto de legendagem em vídeos com texto em tela (*on-screen texts*), um conceito bastante conhecido no mercado de legendagem e bastante pertinente ao tema deste trabalho. Os textos em tela podem ser entendidos como narrativas que podem ser classificadas tanto como elementos de natureza visual verbal não-sonora – quando há apenas a apresentação imagética e nenhuma narração está presente – quanto como visual verbal não-sonora – quando há a narração, normalmente feita pela personagem/falante em cena ou um narrador. No caso de textos em tela classificados como visuais verbais sonoros – mais comuns de ocorrerem em vídeos de cunho educacional ou enfático, como vídeos de marketing –, pode haver a sobreposição de tal classificação, e o mesmo conteúdo verbo-visual é

narrado através do canal verbal sonoro. Exemplos podem ser encontrados em programas cuja temática seja um quiz, e as perguntas verbalizadas pelo apresentador também aparecem escritas na tela.

Em minha experiência enquanto legendista, ora recebia instruções para manter as legendas nos segmentos cujos textos em tela eram idênticos ao texto de legenda, ora era instruído a deletar tais legendas para não haver informações redundantes. Afinal, o texto em tela, visto a partir da exibição de um vídeo, pode ou não funcionar como uma legenda pré-existente? Neste caso, qual seria a estratégia ideal de abordagem tendo em vista o conforto na leitura por parte do público-alvo: repetir a mesma informação textual no texto de legenda ou adequá-la de modo a não haver competição redundante sobre a mesma informação?

Independentemente do tipo de estratégia utilizada, as legendas em produções audiovisuais são usadas amplamente não apenas em produções cinematográficas de grande alcance em vários idiomas, através de legendas interlinguísticas, mas também em outros meios de comunicação, como redes sociais. Nessas plataformas, em especial, o uso de textos em tela é bastante comum e, muitas vezes, esses elementos competem com a informação de legenda que é apresentada, conforme mencionado. Dessa forma, o presente trabalho busca explorar, de maneira inédita, como essas dinâmicas entre textos em tela e legendas influenciam a experiência visual do espectador.

Para discorrer sobre essa temática, o presente trabalho está dividido em 6 capítulos principais, incluindo a Introdução e as Considerações Finais, buscando oferecer não apenas *insights* sobre o tema, mas também uma proposta metodológica que pode ser replicada em estudos futuros.

Começando pelo Capítulo 2, a Fundamentação Teórica, na seção 2.1, “A Legendagem Tradicional na Tradução Audiovisual”, busca-se localizar e conceituar a legendagem dentro dos Estudos da Tradução, apresentando suas classificações e parâmetros, bem como sua contextualização no que tange à recepção por parte do público e ao uso de legendas como textos em tela.

Na seção 2.2, são apresentadas as definições do que são textos em tela, bem como um breve compilado de vídeos com texto em tela, tanto em filmes disponíveis em plataformas de *streaming* quanto em vídeos presentes em redes sociais, tais quais YouTube e Instagram. Na seção, propõe-se a descrição de exemplos e estilo de legendas usadas na presença de textos em tela com

informações semelhantes e, com base nisso, uma classificação que rege as condições experimentais desta pesquisa, a saber: legendas totais e parciais, de acordo com os critérios de estilo observados nos vídeos.

As discussões dessa seção no capítulo ajudam a entender as perguntas iniciais que colaboram para a natureza piloto exploratória da investigação. Isso porque, ao mesmo tempo em que se acredita que o uso de legendas em fragmentos do vídeo que apresentam a mesma informação textual em tela pode levar o espectador a dividir sua atenção, a ausência delas também poderia ocasionar uma redistribuição do olhar entre diferentes partes da tela.

Dessa forma, a problemática deste trabalho pode ser sintetizada na seguinte pergunta central: “Qual é o design de legendas mais eficaz na presença de textos em tela, considerando a experiência visual e a organização da atenção do espectador, especialmente em contextos de redundância informacional?”

A partir dessa questão principal, desdobram-se perguntas específicas que norteiam os objetivos e a metodologia desta pesquisa, como:

1. Como legendas totais e parciais influenciam a forma como o espectador distribui sua atenção entre as áreas de interesse (legenda e texto em tela)?
2. De que forma a redundância informacional nas legendas totais impacta a compreensão do espectador?
3. É possível propor um modelo de design de legendas que facilite a distribuição visual da informação e melhore a experiência de leitura em vídeos com textos em tela?

Para cada uma dessas perguntas, formulam-se as hipóteses que embasam o experimento:

- Hipótese 1: Legendas totais levarão a uma maior dispersão visual e poderão dificultar o processamento da informação devido à redundância entre a legenda e o texto em tela.
- Hipótese 2: Legendas parciais, ao evitarem redundância, promoverão uma distribuição mais fluida da atenção visual e maior clareza na recepção da informação.
- Hipótese 3: A partir dos dados obtidos, será possível propor diretrizes para um design de legendas que otimize a apresentação das informações, equilibrando a experiência visual e textual.

Para entender melhor esses conceitos e buscar responder a essas perguntas, na seção 2.3, são apresentadas as principais pesquisas encontradas com foco investigativo próximo ao tema deste trabalho, no que tange ao estudo da movimentação ocular em vídeos legendados e com textos em tela, embora não concomitantes. Ainda assim, até onde se pôde observar, são escassas as pesquisas com foco investigativo nesses estilos de legenda ao se comparar com pesquisas envolvendo o uso de sistemas de rastreamento ocular e legendagem em geral. Uma possível explicação pode se dar pelo fato de que, para além de as legendas dinâmicas com texto em tela terem um caráter relativamente recente quanto ao seu uso e circulação nas redes sociais, a falta de diretrizes e o caráter não-padronizado da produção também dificulta a elaboração e aplicação de metodologias de estudos, de forma a se obterem variáveis controláveis a serem estudadas.

Além disso, a própria terminologia não consensual decorrente da natureza plural desses textos em tela também dificulta a documentação de trabalhos relacionados à área por palavras-chaves, por exemplo. Nornes (1999 apud Caffrey, 2009, p. 10) utiliza o termo “legendas abusivas” ou “legendagem abusiva” para descrever o uso de legendas que fogem do padrão convencional, mais especificamente “o uso de procedimentos para a criação de legendas que são experimentais, tanto gráfica quanto linguisticamente”. Künzl e Ehrensberger-Dow (2011) utilizam a terminologia “legendas inovadoras” (*innovative subtitles*). O’Hagan e Sasamoto (2016), dentre as pesquisas apresentadas aqui, abordam o uso de “legendas de impacto”, e, em todos esses casos, Sasamoto (2024) as compila dentro de um grande guarda-chuva como “textos em tela”, termo também preferível aqui para descrever elementos verbo-visuais adicionados ao vídeo ou inerentes a ele, como informação que está além da legenda convencional também adicionada ao vídeo.

Essas três seções principais resumem o Capítulo 2 e, com isso em mente, segue-se para o Capítulo 3, Desenho Experimental, dividido em 8 seções, que abordam as variáveis, tipo de pesquisa, bem como os critérios adotados, tendo em vista a influência que questões de natureza linguística e técnica podem ter nos resultados em pesquisas com rastreamento ocular. Sendo assim, na preparação do material, optou-se por se gerar vídeos controlados para evitar variáveis incontroláveis. Esse controle permite uma análise mais precisa das

interações dos participantes com as legendas, contribuindo para um melhor entendimento sobre os padrões de leitura em vídeos legendados, considerando-se os dois tipos de legendas apresentados no Capítulo 2.

Na sequência, nos Capítulos 4 e Capítulo 5, apresentam-se os resultados e discussões suscitados a partir das medidas de análise descritas na seção da metodologia, tais quais: (1) tempo de permanência nas áreas de interesse, (2) número e (3) duração média das fixações, (4) análise de sequência de fixações e (5) questionário pós-coleta.

Como objetivo geral, pretende-se investigar, de forma exploratória, como o design de legendas em vídeos com textos em tela influencia o processamento de leitura e o comportamento visual de espectadores, propondo uma metodologia inicial replicável e adaptável a outros contextos de pesquisa na tradução audiovisual.

Os objetivos específicos incluem (1) analisar a influência da exclusão e permanência de legendas na atenção e na compreensão do espectador em contextos com redundância de informações verbo-visuais, contribuindo para entender o comportamento do público; (2) examinar padrões de fixação e a sequência de movimentos oculares durante a visualização de vídeos legendados com textos em tela; (3) avaliar se a metodologia empregada pode ser válida para pesquisas nesse contexto.

Sendo assim, em suma, este estudo apresenta-se como uma pesquisa inicial no uso de rastreamento ocular para analisar a interação entre textos em tela e legendas em vídeos voltados para redes sociais. Dada a escassez de estudos que abordem esse tipo de material audiovisual, o trabalho busca propor uma metodologia pioneira que possa ser replicada em contextos futuros, contribuindo para o desenvolvimento de modelos de legenda e melhores práticas na área de tradução audiovisual.

Ao explorar o impacto dos textos em tela na dinâmica de leitura de legendas, este estudo visa preencher uma lacuna existente nos estudos de tradução audiovisual, especialmente no que tange ao uso de ferramentas tecnológicas, como o rastreamento ocular, para avaliar o comportamento visual de espectadores em plataformas digitais. A ausência de padronização nas diretrizes metodológicas dessa área ressalta a relevância de se propor um modelo experimental que possa servir de base para investigações futuras.

Além de abordar aspectos teóricos sobre a interação entre texto e imagem no campo audiovisual, esta pesquisa enfatiza a importância de compreender como diferentes condições de legendagem influenciam o foco visual. Dessa forma, este trabalho não apenas oferece insights preliminares sobre o tema, mas também propõe um marco metodológico para ampliar as possibilidades de análise e aplicação prática na área de tradução audiovisual.

## 2 FUNDAMENTAÇÃO TEÓRICA

Neste capítulo, abordam-se as principais discussões teóricas que embasam os objetivos e a metodologia deste trabalho. Está dividido em três seções, que apresentam as definições dos termos, o contexto de pesquisa e trabalhos relacionados ao uso do rastreamento ocular em legendagem.

### 2.1 A Legendagem Tradicional na Tradução Audiovisual

A Tradução Audiovisual (TAV) tornou-se, nas últimas décadas, umas das áreas mais prolíficas dentro dos Estudos da Tradução, apesar de ter sido ignorada por anos como um campo de estudo cujo foco investigativo merecesse atenção (Díaz Cintas; Remael, 2021). Segundo Díaz Cintas e Remael (2021), por vários anos, a TAV não foi vista como uma área de pesquisa dentro dos Estudos da Tradução, dado que ela enfrenta limitações de natureza temporal e espacial para ser realizada, o que interfere diretamente no resultado; ou seja, há limitações em relação ao número de caracteres por linha, bem como ao tempo disponível de exibição em que se pode usar tais caracteres no que diz respeito à legendagem, por exemplo. Mesmo os Estudos da Tradução só passaram a se consolidar como campo disciplinar independente a partir da década de 1980 (Munday, 2001).

O próprio título “Tradução Audiovisual” também só passou a ser usado no final dos anos 1980s, a partir da consolidação dos Estudos da Tradução e da popularização do VHS - *Video Home System* (Gonçalves, 2021). Isso porque, com a popularização do VHS entre as décadas de 1980 e 1990, as produtoras de filmes começaram a se dedicar à criação de legendas, algo que antes era feito apenas para filmes exibidos nos cinemas (Franco; Araújo, 2012; Nuñez, 2017). Até então, os trabalhos voltados à TAV eram comumente mencionados como *film translation* (Tradução de filmes) (Franco; Araújo, 2012). Díaz Cintas e Remael (2021) complementam, apontando que a TAV, como prática profissional, já existia desde a invenção do cinema, na virada do século XX, mas somente a partir da metade dos anos 1990s que ganhou mais adeptos. Araújo (2002) traz ainda que as pesquisas nessa área passaram a acontecer, principalmente, a partir dos anos 1980s, na Europa, e a partir dos anos 1990s, no Brasil.

Dentro de sua categoria, apontada por Díaz Cintas e Remael (2021) como um termo “guarda-chuva”, a TAV, de forma geral, caracteriza-se como uma compilação de várias práticas de tradução no campo audiovisual, que diferem no que tange ao resultado linguístico e às estratégias utilizadas, dentre elas, interpretação, dublagem e legendagem. E, mesmo dentro dos grandes termos, a TAV apresenta suas submodalidades.

### **2.1.1 Classificação**

Gambier (2003) busca apresentar uma visão geral das modalidades, que Franco e Araújo (2012) apontam como confusa, a saber: legendagem interlinguística ou legenda aberta (*interlingual subtitling* ou *open caption*), legendagem bilíngue (*bilingual subtitling*), dublagem (*dubbing*), dublagem intralingual (*intralingual dubbing*), interpretação consecutiva (*consecutive interpreting*), interpretação simultânea (*simultaneous interpreting*), interpretação de sinais (*sign language interpreting*), voice-over ou meia-dublagem (*voice over* ou *half dubbing*), comentário livre (*free commentary*), tradução à prima vista ou simultânea (*simultaneous or sight translation*), produção multilinguística (*multilingual production*), legendagem intralinguística ou *closed caption* (*intralingual subtitling* ou *closed caption*), tradução de roteiro (*scenario/script translation*), legendagem ao vivo ou em tempo real (*live or real time subtitling*), supra-legendagem ou legendagem eletrônica (*surtitling*) e audiodescrição (*audiodescription*). As autoras acrescentam que, apesar da grande variedade citada por Gambier (2003), muitas podem acontecer no meio audiovisual, mas não necessariamente fazerem parte dele. Na sequência, afunilam a discussão apresentando as principais modalidades em três tópicos, a saber: legendagem, modos de revocalização (dublagem e voice-over) e audiodescrição. Díaz Cintas e Remael (2021) também apresentam uma proposta de classificação das várias atividades que compõem a TAV e utilizam o termo “guarda-chuva” (*umbrella*) para se referirem à área, já que esta engloba tais atividades, incluindo a legendagem, classificada de forma geral como a prática de adicionar um texto escrito a uma produção original, a saber, *Timed Text* (texto temporizado).

No que concerne à Legendagem, uma vez que faz parte do escopo deste trabalho, levando também em consideração os avanços ocorridos nos últimos

anos, há autores que se detêm a duas categorias gerais. São eles Araújo (2022) e Díaz Cintas e Remael (2021), que classificam os tipos de legenda segundo os parâmetros técnicos e linguísticos.

No que toca o parâmetro linguístico, as legendas podem ser classificadas como intralinguísticas, quando são feitas dentro do mesmo idioma, ou interlinguística, as mais conhecidas, quando traduzidas de um idioma a outro. Há também legendas que podem ser exibidas em dois idiomas ao mesmo tempo, as legendas bilíngues (Díaz Cintas; Remael, 2021). Sendo assim, a legendagem - tal qual a definição que aborda a tradução como um todo, também trazida por Jakobson (1995), e como parte dos Estudos da Tradução - pode ser classificada como interlinguística e intralinguística. Contudo, alguns estudiosos relutam em considerar a legendagem intralinguística como tradução (Díaz Cintas; Remael, 2021), devido às restrições de natureza temporal e espacial, mencionadas anteriormente. Nessa perspectiva, as alterações e reformulações textuais a “distanciaria” do texto de partida.

No que diz respeito aos parâmetros técnicos, as legendas podem ser classificadas como abertas ou fechadas (Araújo, 2002; Díaz Cintas; Remael, 2007, 2021). No primeiro caso, as legendas são apresentadas no vídeo de forma fixa, permanente, ou seja, não podem ser ativadas ou desativadas pelo espectador durante a transmissão. No segundo, opostamente, as legendas podem ser ativadas ou desativadas, dependendo das preferências do espectador, como em algumas plataformas de streaming, que não apenas permitem ao usuário escolher a legenda em um dado idioma, como também alterar sua cor, tamanho e estilo.

Em ambas as categorias, as legendas podem ser intra ou interlinguísticas, e as classificações apresentadas aqui não necessariamente anulam ou impedem outras classificações e subclassificações. Assim, uma legenda aberta pode ser intra ou interlinguística, assim como uma legenda fechada pode também ser classificada como tal. Pode haver ainda a possibilidade de um mesmo produto audiovisual apresentar legendas abertas e fechadas, quando há, por exemplo, outro idioma na trilha, que é mostrado através do que é conhecido como “*forced narratives*” (FN), legendas que visam fornecer informações necessárias para a compreensão dos espectadores (Díaz Cintas; Remael, 2021).

Além dos parâmetros técnicos e linguísticos apresentados, Díaz Cintas e

Remael (2021) se propõem também a expandir as discussões sobre legendagem de acordo com alguns aspectos, a saber: (1) o tempo disponível para preparação, (2) os modos de exibição, (3) métodos de projeção e (4) meios de distribuição. Cada aspecto é brevemente discutido a seguir, por princípios de contextualização.

#### 2.1.1.1 *Tempo disponível para preparação*

No que tange ao primeiro aspecto, os autores apresentam os tipos de legendas compreendidas quanto ao tempo disponível para preparação: Legendas pré- preparadas/pré-gravadas, como dedutível a partir da própria terminologia, são produzidas a partir do material de vídeo já gravado. Isso permite que os tradutores legendistas trabalhem com tempo prévio em cima do material audiovisual produzido, com precisão de sincronia entre som e imagem, tempo de pesquisa, além de aplicação das regras e diretrizes, que são conhecidas no mercado de legendagem de acordo com cada distribuidora ou streaming, tais quais CPS (caracteres por segundo) e CPL (caracteres por linha) adequados, número de linhas, segmentação linguística, retórica e visual – ver Reid (1990) –, bem como controle de qualidade e correções de eventuais erros. Tais características compõem o tipo de legenda mais conhecido.

No caso das legendas ao vivo, semi ao vivo ou em tempo real, esse tempo de preparação pode não acontecer, dando margem a eventuais erros ou inconsistências, tanto em relação a questões técnicas quanto linguísticas. Podem ser geradas de forma automática, por reverberação ou através do trabalho do estenotipista. Em resumo, em qualquer dessas situações, as legendas podem ser editadas para comportar as restrições de tempo e espaço, ou *verbatim*.

#### 2.1.1.2 *Modos de exibição*

Em relação ao modo de exibição, ainda segundo Díaz Cintas e Remael (2021), as legendas podem aparecer de forma completa em blocos, *pop-ups* ou *pop-nos* – como o público tende a estar mais habituado –, ou de forma cumulativa, quando necessário atrasar o surgimento de dada legenda em

contexto pertinente – como em pausas dramáticas, por exemplo (William, 2009). Outras estratégias, como dividir o texto em duas legendas, também são comuns nos *streamings*, como na Netflix – ver Netflix, (N.D). Já as legendas no estilo *roll-up*, ou rotativas, aparecem com fluxo de texto constante, palavra por palavra, ou em fases curtas. Esse estilo de legenda costuma ser comum em *closed captions* e pode aparecer em até quatro linhas em programas ao vivo (Araújo, 2002). “Tendem a não ser apreciadas pelos espectadores, porque sua instabilidade e o constante movimento ascendente do texto em tela torna a leitura uma atividade cognitiva bastante desafiadora e onerosa (Díaz Cintas; Remael, 2021, p. 26, tradução nossa<sup>1</sup>)”.

### 2.1.1.3 Métodos de projeção e meios de distribuição

Ambas as classificações trazidas por Díaz Cintas e Remael (2021) estão ligadas, como pontuam os autores, à evolução tecnológica pela qual transitou o processo de criação e veiculação de legendas ao longo dos anos. De um lado, apresentam-se os mecanismos usados para exibição, tais quais: legendagem termal ou mecânica, fotoquímica, ótica, a laser, eletrônica, 3D ou imersiva; do outro lado, apresenta-se onde essa exibição será feita, a saber: cinema, vídeo, VHS, DVD, VCD, *Blue-ray*, televisão ou internet.

Empiricamente, é correto dizer que, ao passo que processos de legendagem no passado eram comuns para o uso em VHS, por exemplo, atualmente, esses processos abrem vaga para outras mídias. Sendo assim, as legendas ainda são bastante usadas no cinema e televisão, mas possuem grande impacto e presença também na internet, com a popularização das plataformas de streaming.

Com tantos processos e classificações, é de se esperar que materiais legendados passem por “avaliação” do público que o consome, que, em sua grande maioria, desconhece tais processos e aplicações de acordo com cada objetivo, método ou veículo de divulgação. Embora não seja o foco deste trabalho, considera-se relevante abordar brevemente este tópico na seção seguinte, especialmente no que se refere ao senso comum sobre a legendagem,

---

<sup>1</sup>They tend to be disliked by viewers because their instability and constant upward movement of the text on screen make the reading a rather challenging and onerous cognitive activity.

dado que o público-alvo deste trabalho é a grande massa consumidora de materiais legendado.

### ***2.1.2 A Legendagem para o público em geral***

A tradução, seja ela de obras literárias ou de conteúdo audiovisual, vai passar por críticas, sob os olhares daqueles que têm o acesso e a chance de consumir o produto apresentado em sua língua de partida. O nível dessa crítica e o tempo até que ela ganhe repercussão dependerá do *gap* entre a publicação da obra de origem e de sua tradução.

Esse tema é apresentado por Díaz Cintas e Remael (2007), ao pontuarem as constantes críticas, geralmente negativas, direcionadas às legendas de produtos audiovisuais destinados ao público em geral, como filmes, séries e documentários televisivos. Isso, porque, segundo os autores, o produto audiovisual legendado permite ao público ter acesso simultâneo a ambas as línguas, a de chegada e a de partida, gerando esse efeito imediato de crítica. Ou seja, quando um espectador reconhece uma palavra na língua estrangeira e percebe que a tradução desta pode não ter sido feita da forma como a reconhece, gera-se esse efeito de estranhamento e comentários do tipo “não foi isso que foi dito” ou “essa tradução está errada”, quando na verdade, o processo de legendagem está sujeito a limitações de tempo e espaço (e.g., CPL e CPS), que requerem do tradutor adaptações em relação ao texto de partida, seja esta ao nível de palavra ou oração.

No que tange ao CPS, os valores limites podem variar entre distribuidoras e clientes, mas, em geral, se mantêm na faixa de 14 a 25 caracteres por segundo, dependendo do público-alvo e idioma. Teóricos como Díaz Cintas e Remael (2007) também propõem valores medidos em “palavras por segundo”, com três velocidades consideradas confortáveis para leitura, a saber: 145 ppm (palavras por minuto); em torno de 160 ppm; em torno de 180 ppm. Autores como Nascimento (2018), Assis (2021), Vieira (2016), e Naves, *et al* (2016) consideram essas velocidades respectivamente como lenta, média e rápida.

Díaz Cintas e Remael (2021) sugerem alguns valores de conversão entre ppm e cps. Os valores são expostos abaixo, no Quadro 1, incluindo espaços em branco, e no Quadro 2, excluindo espaços em branco.

**Quadro 1** – Conversão de valores entre palavras por minuto (ppm) e caracteres por segundo incluindo espaços em branco.

| 12 CPS  | 13 CPS  | 14 CPS  | 15 CPS  | 16 CPS  | 17 CPS  | 18 CPS  | 19 CPS  | 20 CPS  |
|---------|---------|---------|---------|---------|---------|---------|---------|---------|
| 150 ppm | 160 ppm | 170 ppm | 180 ppm | 190 ppm | 200 ppm | 215 ppm | 225 ppm | 240 ppm |

Fonte: traduzido de Díaz Cintas e Remael (2021, p. 112).

**Quadro 2** – Conversão de valores entre palavras por minuto (ppm) e caracteres por segundo excluindo espaços em branco.

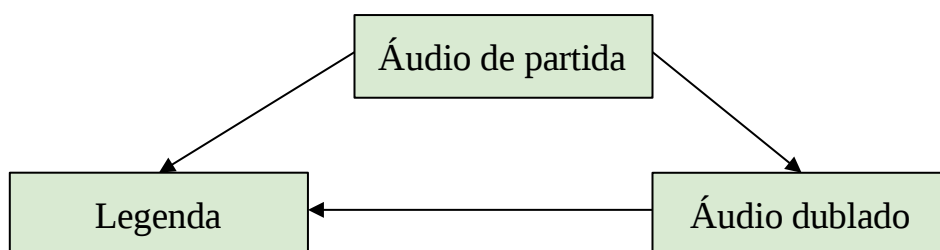
| 12 cps  | 13 cps  | 14 cps  | 15 cps  | 16 cps  | 17 cps  | 18 cps  | 19 cps  | 20 cps  |
|---------|---------|---------|---------|---------|---------|---------|---------|---------|
| 130 ppm | 140 ppm | 150 ppm | 160 ppm | 170 ppm | 180 ppm | 190 ppm | 200 ppm | 215 ppm |

Fonte: traduzido de Díaz Cintas e Remael (2021, p. 112).

O CPS é um dos parâmetros técnicos da legendagem, os quais são necessários de serem seguidos, além dos parâmetros linguísticos. Os aspectos técnicos abrangem elementos como quantidade de linhas, limite de caracteres por linha, escolha de cores, formato, posição na tela, duração de exibição e velocidade de leitura (Assis, 2016, p. 28). Já os parâmetros linguísticos envolvem a segmentação adequada do texto e ajustes no conteúdo textual, como cortes, reduções, reestruturações ou condensações de informações (ASSIS, 2016, p. 26). Sendo assim, como pontuam Díaz Cintas e Anderman (2008), o legendista não tem espaço para formulações robustas com estruturas frasais complexas e precisa ser sucinto para melhorar a legibilidade.

Outro exemplo de críticas pode surgir a partir do que considero aqui como um “triângulo analítico de tradução”. Isto é, quando esse tipo de crítica é feito em cima de um produto audiovisual dublado, comparando-o ao áudio de partida, bem como à sua legenda na língua de chegada, tanto em relação ao áudio de partida quanto ao áudio dublado. Sendo assim, a crítica destina-se não somente à comparação com o áudio de partida como também ao áudio dublado.

**Figura 1** – Esquema ilustrativo de produção de legendas a partir de duas fontes



Fonte: elaborado pelo autor.

No esquema acima, são apresentados dois processos tradutórios que se encontram dentro do mesmo nicho na Tradução Audiovisual, mas que seguem estratégias diferentes. O mesmo produto audiovisual, seja filme ou séries – mencionados aqui por serem o tipo de conteúdo mais consumido mundialmente –, não necessariamente é traduzido para dublagem e legendagem pela mesma pessoa; ou seja, um legendista pode trabalhar de forma individual em um dado episódio, enquanto outra pessoa faz/fará a tradução para dublagem.

Ainda nessa linha de pensamento, pode surgir a ideia equivocada de que a legenda deve seguir o áudio dublado, gerando outro tipo de crítica: a de que a legenda não é “fiel” ao áudio apresentado. No entanto, esse processo não ocorre necessariamente dessa forma, devido a três principais motivos. (1) Dublagem e legendagem são processos diferentes de tradução e, portanto, seguem estratégias e técnicas diferentes. (2) A legendagem geralmente é produzida a partir do áudio original do conteúdo, e não com base no áudio dublado, podendo haver situações em que se pede ao legendista para utilizar o roteiro de dublagem ou o áudio dublado, ao invés do áudio ou roteiro de partida. (3) Além disso, mesmo que o legendista use o áudio dublado, ainda assim a legenda não estará 100% semelhante ao que é dito, já que adaptações de caráter técnico e linguístico costumam ser feitas na legendagem para que ela se adeque, principalmente, ao CPS e CPL – como a não repetição de palavras ou a omissão de vícios de linguagem; exceto em casos em que é pedido um arquivo de *closed captions* (CC). Mesmo assim, são casos específicos que ocorrem

apenas quando solicitados por parte do cliente ou plataforma na qual o produto será veiculado, como a Netflix (2024)<sup>2</sup>.

A título de exemplo, no que diz respeito às repetições, em seu guia para a produção de legendas, a Netflix define: “Não traduzir palavras ou frases repetidas mais de uma vez pelo mesmo falante”, e complementa: “Se a palavra ou frase repetida for dita duas vezes seguidas, sincronizar a legenda com o áudio, mas traduzir apenas uma vez”<sup>3</sup>. Já com relação ao uso do script de dublagem em português, a plataforma indica que

Quando o áudio for dublado em português brasileiro, por favor, usar o áudio ou o roteiro dublado como base para a produção do arquivo LSE, a fim de garantir que ambos estejam o mais semelhante possível, desde que a velocidade de leitura e sincronia permitam (tradução nossa)<sup>4</sup>.

### **2.1.3 Por que o uso de legendas intralinguísticas?**

Sendo considerada um tipo de tradução audiovisual, a legendagem pode ser uma forma de tradução intralinguística, interlinguística ou intersemiótica (Araújo; Assis, 2014). Esses conceitos também são discutidos por Jakobson (1995) e são frequentemente utilizados dentro dos Estudos da Tradução. O primeiro caso ocorre dentro de um mesmo sistema linguístico, ou seja, quando o idioma do áudio é o mesmo da legenda; já o segundo pode ser associado à legendagem de filmes estrangeiros, em que o idioma do áudio e da legenda são diferentes. Analisando a legendagem ainda em uma perspectiva mais ampla, incluindo aqui a LSE, Legendagem para Surdos e Ensurdidos, pode-se também entender esta como uma forma de tradução intersemiótica (Assis, 2016), uma vez que, nesse tipo específico de legenda, é feita a tradução dos efeitos sonoros (Araújo; Assis, 2014), entendidos como signos de natureza não verbal (Jakobson, 1995). Tendo em mente que, nas redes sociais, o público tem a opção de navegar no *feed* com o áudio ativado ou desativado, é de se esperar que parte da população

<sup>2</sup>Disponível em: <<https://partnerhelp.netflixstudios.com/hc/en-us/articles/215600497-Brazilian-Portuguese-Timed-Text-Style-Guide>>. Acesso em: 19 de Ago. 2022.

<sup>3</sup> *If the repeated word or phrase is said twice in a row, time subtitle to the audio, but translate only once. Do not translate words or phrases repeated more than once by the same speaker* Tradução nossa).

<sup>4</sup> *Where content has been dubbed into Brazilian Portuguese, please refer to the dubbed audio or dubbing script as the basis for the SDH file and ensure that the two match as much as reading speed and timings allow.*

consumidora de tal conteúdo tenha o hábito de navegar com o áudio desativado. De fato, segundo Tirumala e Youngblood (2021), 85% dos usuários navegam nas redes sociais com o áudio desativado. Infere-se, assim, que, se o conteúdo em vídeo não lhes for atrativo apenas através das imagens, vídeos cujo conteúdo poderia lhes despertar certo interesse poderiam passar facilmente despercebidos no *feed*, não fosse o uso de legendas, que podem atrair a atenção do usuário.

A partir desse ponto, o uso de legendas intralinguísticas em vídeos produzidos para redes sociais assume também um caráter tanto inclusivo quanto mercadológico. Inclusivo, na medida em que promove ao público surdo e ensurdecido acesso à informação sonora; mercadológico, porque busca atingir um número maior de consumidores, sejam estas pessoas surdas ou pessoas ouvintes que se encontram dentro dos 85% de usuários que navegam nas redes sociais com o áudio desativado.

## **2.2 On-Screen Texts (Textos em Tela)**

Esta seção destina-se a apresentar brevemente exemplos dos tipos de legenda que se pretende observar neste trabalho, de acordo com suas formas de apresentação em plataformas sociais, como YouTube e Instagram. Não obstante, também se faz necessário entender o que são textos em tela.

Dentre as pesquisas analisadas para esta dissertação, foram encontrados poucos trabalhos com foco no tema de textos em tela, principalmente no quesito definição e uso em redes sociais, que é o foco deste trabalho. Pesquisas voltadas à análise desse segmento de produções audiovisuais ainda parecem escassas até onde se pôde observar. Elas apresentam-se com foco predominante na análise de programas de TV japoneses, em especial noticiários, como os trabalhos de Sasamoto (2024), Sasamoto, O'Hagan e Doherty (2017), O'Hagan (2010), Sasamoto, Doherty e O'Hagan (2021), O'Hagan e Sasamoto (2016) e Park (2008; 2009).

O uso de textos em tela se originou no início do século XX, com o uso de intertítulos no cinema mudo (Sasamoto, 2024). Os intertítulos também marcam o início das legendas e “são conhecidos como cartões de título, podendo ser definidos como um pedaço de texto impresso, filmado, que aparece entre as

cenar (Díaz Cintas; Remael, 2021, p. 30<sup>5</sup>). A principal função dos intertítulos no cinema mudo era representar os diálogos dos personagens e os materiais de narração descritivos relacionados às imagens.

Com o passar dos anos e com o advento de legendas interlinguísticas, ou seja, ao se traduzir o conteúdo audiovisual, no entanto, essas informações precisam ser passadas para o público dentro do mesmo arquivo de legenda, ou em um arquivo de legenda à parte, através da inserção das *Forced Narratives* (Díaz Cintas; Remael, 2021). O último caso costuma acontecer na exibição de obras dubladas, quando tais trechos não foram dublados para o idioma de chegada, tais quais placas, letreiros, mensagens etc. *Forced Narratives*

são uma sobreposição de texto que esclarece as comunicações ou idiomas alternativos que devem ser compreendidos pelo espectador. Também podem ser usadas para esclarecer diálogos, gráficos de texto ou IDs de localização/pessoa que não são abordados no áudio dublado/localizado. Para permitir a mesma experiência de visualização em vários países e dispositivos, as legendas FN são localizadas e entregues como arquivos de texto cronometrados separados (NETFLIX, 2017, tradução nossa).

Arquivos de legendas FN podem ser usados para incluir a tradução de um diálogo estrangeiro, transcrição de diálogos inaudíveis que tenham sido incluídos na transmissão ou apresentação teatral, ou tradução de textos em tela que sejam relevantes para o entendimento do contexto, como letreiros, textos criativos etc. (Netflix, 2023), como exemplificado na imagem abaixo:

---

<sup>5</sup> They are also known as title cards and can be defined as a piece of filmed, printed text that appears between scenes.

**Figura 2** – Texto em tela: ID Gráfico de localização traduzido com *Forced Narratives*



Fonte: Netflix (2017).

Mais exemplos de textos em tela são apresentados na seção 3 deste trabalho. Em alguns casos, como apontam Díaz Cintas e Remael (2021), os intertítulos são traduzidos através da dublagem ou voice-over ou mesmo na forma de narração, sendo mantidos os gráficos de origem. Em exemplos mais elaborados, com casos que envolvem edição de vídeo, tais informações visuais são visualmente editadas para compreender o público-alvo. Exemplos deste tipo podem ser encontrados em filmes, como “Divertidamente” (Lins, 2022). De todo modo, tanto as *forced narratives* quanto os intertítulos, no momento de seu advento, podem ser considerados como exemplos de texto em tela.

Retomando a discussão sobre o uso terminológico, segundo Sasamoto (2024), fora do Japão, não há uma terminologia definida para o que se conhece como “*text on screen*” (texto em tela); o termo “*caption*” ou legenda parece prevalecer, principalmente na medida em que os textos em tela se tornam cada vez mais frequentes nas redes sociais. A fim de facilitar a compreensão, *on-screen texts* (textos em tela), no que tange aos objetivos deste trabalho, são informações verbo-visuais inter ou intralinguísticas presentes em vídeo, as quais são relevantes para o entendimento do material audiovisual, podendo ou não ser narradas.

Esta é uma definição geral do que se costuma ver em obras audiovisuais em geral. Sasamoto, O’Hagan e Doherty (2017) também contribuem para essa definição ao trazer para a discussão outros exemplos de textos em tela,

apresentados como telop (*television opaque projector*), amplamente usados em programas de TV japoneses. Segundo os autores,

Pesquisas sobre TV e novas mídias publicadas em inglês ignoraram técnicas de edição de pós-produção que adicionam legendas no mesmo idioma usado fora do mundo anglófono. Um caso em questão é uma forma única de legendas de TV que tem sido um recurso regular da TV japonesa há algum tempo<sup>6</sup> (Sasamoto, O'hagan, Doherty, 2017, p. 1, tradução nossa).

Os autores complementam, citando O'Hagan (2010), que produtores televisivos geralmente usam o telop para atrair a atenção dos espectadores a um dado elemento em uma cena específica de um programa e que essas legendas intralinguísticas podem ser consideradas um componente midiático que vai além de apenas transmitir uma transcrição do que é falado (Sasamoto, O'hagan, Doherty, 2017). Esse estilo de legenda é definido por Park (2008 *apud* Sasamoto; Doherty; O'hagan, 2021) como legendas de impacto (*impact captions*), e possuem as seguintes características :

Contém conteúdo literal, parafraseado e editorializado com uma camada extra de informações e significado; podem aparecer em qualquer lugar da tela; são ricas em informações sensoriais multimodais; são usadas pelos produtores para obter um efeito específico na recepção por parte espectador (por exemplo, humor, surpresa, ênfase, incitamento ao preconceito, melhor compreensão e design narrativo); não podem ser desligadas, ou seja, são 'abertas' (Sasamoto; Doherty; O'hagan, 2017, p. 253, tradução nossa).

Embora não sejam o foco deste estudo, as legendas de impacto através do telop, apresentadas pelos autores acima, ajudam a entender o estilo de legendas dinâmicas que vêm sendo utilizados nas mídias sociais e se tornam relevantes neste trabalho, na medida em que se assemelham às suas características e finalidades. Sendo assim, com o uso cada vez mais frequente das redes sociais, como Instagram, Facebook, YouTube, dentre outras, as definições expostas acima podem ser mais abrangentes, uma vez que produtores de conteúdo nessas plataformas utilizam textos de natureza dinâmica, com legendas e outras informações visuais, geralmente com bastante uso de cores, que captem a atenção do público (Sasamoto, 2024).

---

<sup>6</sup> The research literature on TV and new media published in English has ignored post- production editing techniques of adding the same language captions used outside of the Anglophone world. A case in point is a unique form of TV captions that has been a regular feature of Japanese TV for some time.

Segundo a autora, uma vez que não há diretrizes oficiais para esse estilo de conteúdo, cada produtor o cria de acordo com o que consideram adequado e artístico para seu produto; diferentemente de materiais com diretrizes já estabelecidas, como filmes e séries de grande circulação, ou as legendas de impacto apresentadas acima.

Em suma, apesar de as definições acima estarem voltadas ao uso de textos em tela como informações adicionadas, majoritariamente em materiais intralinguísticos, é importante mencionar que, dentro de uma obra a ser traduzida, informações verbo-visuais diegéticas também podem ser consideradas como tais. Isso se dá, especialmente, pela necessidade de se adicionar uma legenda interlinguística (*forced narratives*) com a tradução de tal informação, caso não haja um processo de localização da obra nem de sua tradução através da narração por dublagem ou voice-over.

No contexto deste trabalho, textos em tela são considerados como qualquer informação verbo-visual que seja parte do vídeo ou tenha sido adicionada a ele em momento anterior à criação de legendas. Com base nisso e no escopo deste trabalho, a seguir, são apresentados alguns exemplos de texto em tela encontrados no *streaming* e em redes sociais, de acordo com as classificações apresentadas aqui.

### **2.2.1 Exemplos de Textos em Tela Interlinguísticos através de Localização Imagética**

Nesta seção, são apresentados alguns exemplos de textos existentes a partir do processo de localização (interlinguístico). No filme “Divertidamente”, da Pixar, há algumas situações nas quais há tradução de texto em tela diretamente na edição do vídeo localizado. Em uma das cenas em questão, a personagem Raiva está lendo um jornal, no minuto 03:47, cuja informação verbo-visual foi traduzida na própria edição de vídeo.

**Figura 3** – Estratégia de Localização em Divertidamente por Edição de Vídeo



Fonte: modificado do filme “Divertidamente”.

Estratégia semelhante pode ser observada também em outros idiomas, como exemplifica Lins (2022):

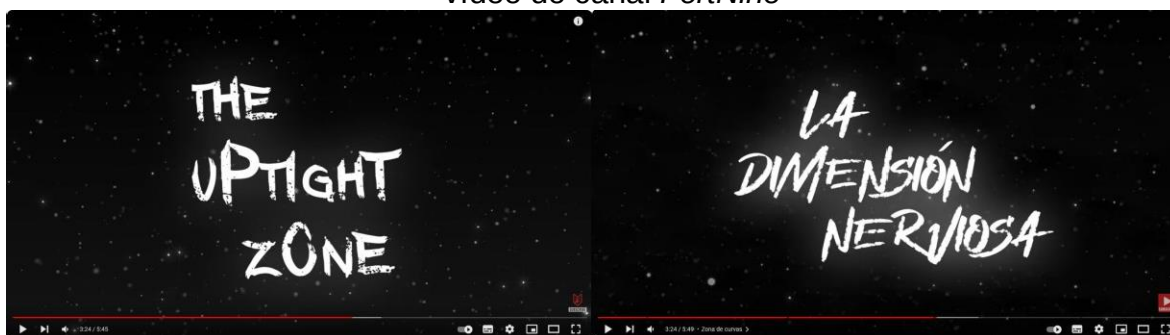
**Figura 4** – Estratégia de Localização 2 em Divertidamente por edição de vídeo



Fonte: Val (2015) e Muniz (2020 *apud* Lins 2022, p. 34).

Situações similares podem ser encontradas a depender do país de recepção e do idioma para o qual o material será localizado. Outros exemplos podem ser encontrados em canais localizados, como pode ser observado abaixo, na localização do vídeo do canal *FortNine* para o espanhol. A imagem à esquerda foi retirada do canal de partida (Fortnine, 2018); na sequência, a localização do intertítulo para espanhol, feita através de edição e animação de vídeo (Fortnine, 2023).

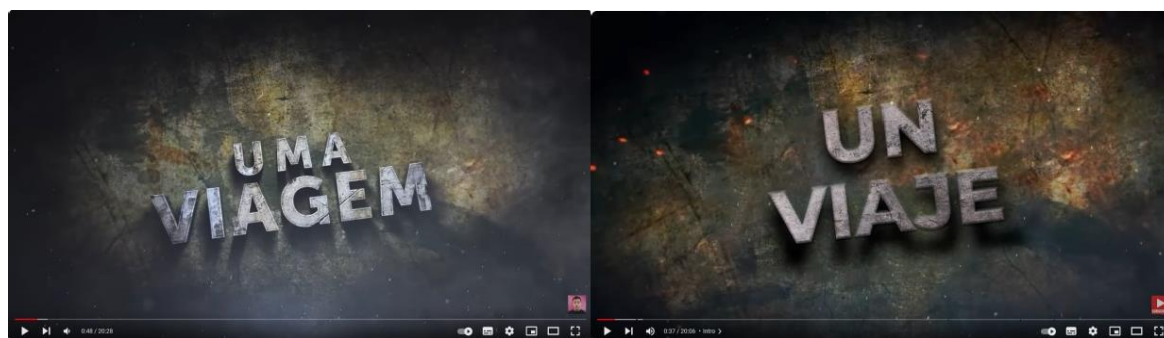
**Figura 5** – Captura de tela: tradução por edição de vídeo para o texto em tela em vídeo do canal *FortNine*



Fonte: FortNine (2018; 2023).

Exemplo similar também pode ser observado no canal do Felipe Neto localizado para o espanhol. A imagem da esquerda mostra o texto “Uma viagem” no canal de partida em português; na imagem da direita, vê-se a localização do intertítulo “Un viaje”, em espanhol.

**Figura 6** – Captura de tela: Tradução por edição de vídeo para o texto em tela em vídeo do canal Felipe Neto



Fonte: Felipe Neto (2023; 2024).

Tais exemplos estão mais relacionados ao processo de localização e não são tão comuns no mercado audiovisual em geral, em decorrência do tempo e custo de produção, bem como da necessidade de haver outro arquivo de vídeo disponível. Tal fato não costuma ser o caso de todos os materiais, já que muitas distribuidoras disponibilizam o mesmo arquivo de vídeo com diferentes faixas de áudio e de legenda para cada idioma.

### 2.2.2 Exemplos de Textos em Tela Interlinguísticos através de *Forced Narratives*

Com base na segunda classificação proposta, apresentam-se aqui alguns exemplos de textos em tela presentes no material audiovisual de origem. Tais exemplos podem ser traduzidos no material de chegada interlinguístico, através da adição de um arquivo de legenda, por meio das *Forced Narratives* ou do próprio arquivo de legenda, quando tal tradução não é fornecida a partir de uma narração por dublagem ou voice-over.

No filme “O Impossível”, por exemplo, o personagem Greg tira do bolso um bilhete, onde está escrito “*We are at the beach*”, que remete a uma cena anterior do filme e, que, necessariamente, precisaria ser traduzida para contextualizar a informação. A tradução por meio das *Forced Narratives*, feita em caixa alta, “ESTAMOS NA PRAIA”, promove ao público essa informação, que seria perdida por quem não fala inglês e, conseqüentemente, poderia prejudicar o entendimento da cena.

**Figura 7** – Captura de tela: texto em tela através de *Forced Narratives* em “O Impossível”



Fonte: Modificado do filme “O Impossível”.

Outros tipos de texto em tela presentes no *streaming* podem ser informações adicionadas, como créditos ou descrições verbais, assim como exemplificado na imagem abaixo.

**Figura 8** – Captura de tela: texto em tela através de *Forced Narratives* em “O Impossível” (2)

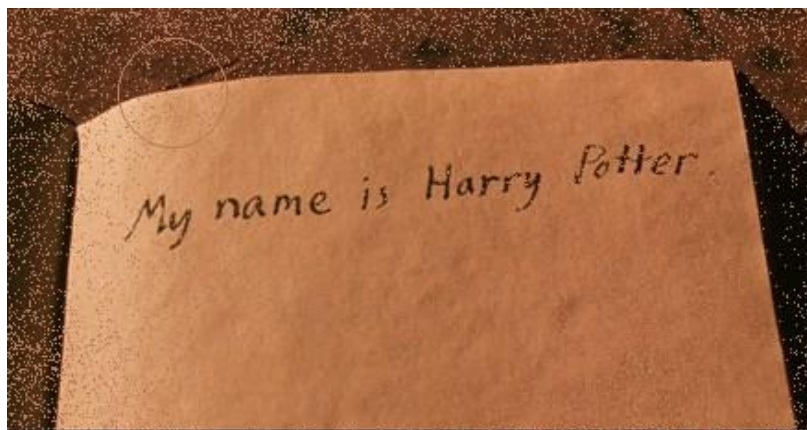


Fonte: Modificado do filme “O Impossível”

Em ambos os casos, há a presença de dois textos em tela. Um primeiro, pertencente ao próprio construto imagético da obra de partida, com base no qual houve a adição de um segundo texto em tela através da legenda no material de chegada.

Essa inserção pode acontecer de duas formas: a primeira, através do próprio arquivo de legenda, caso a presença do texto em tela seja narrada no idioma de partida, como em “Harry Potter e a Câmara Secreta”, em que Harry escreve “*My name is Harry Potter*”, ao mesmo tempo em que fala em voz alta o que está escrevendo; a segunda, através da própria dublagem. Na versão dublada, não há a inserção de um novo texto em tela, já que a própria fala está sendo dublada. Na versão legendada, a tradução “Meu nome é Harry Potter” é feita através do próprio arquivo de legenda, sem a necessidade de inserção de um novo arquivo de legenda como *Forced Narratives*.

**Figura 9** – Captura de tela: texto em tela com narração em “Harry Potter e a Câmara Secreta”



Fonte: Modificado do filme “Harry Potter e a Câmara Secreta”.

Em suma, esses são exemplos em que a tradução dessas informações ocorre quando temos vídeos em idiomas estrangeiros, ou seja, através do uso de legendas interlinguísticas. Vale trazer para a pauta que legendas intralinguísticas são comumente usadas em materiais audiovisuais, que são disponibilizados em diversos estilos, dentre eles, a presença de texto em tela com narração, que pode ocorrer também de maneira interlinguística.

### ***2.2.3 Exemplos de Texto em Tela com narração intralinguística***

Na série “Fã Clube”, da Cartoon Network, encontramos exemplos de um material intralinguístico com bastante informação visual. A série foi produzida de forma localizada, ou seja, foram gravados episódios em outros países sem a necessidade de tradução dos mesmos. O exemplo abaixo foi retirado da versão brasileira da série, disponibilizada na HBO Max e também no canal do YouTube da Cartoon Network Brasil (ver Brasil, 2023). São cinco episódios em estilo Quiz, cujas perguntas são mostradas em tela e faladas ao mesmo tempo pela apresentadora.

**Figura 10** – Captura de tela: Texto em Tela com narração em material intralinguístico



Fonte: Brasil (2023).

No tange à criação de legendas intralinguísticas para materiais desse tipo, como proceder, tendo em vista que grande parte do material possui informação verbo- visual disponível com a mesma narração? Para as diretrizes da HBO Max ou Warner Bros (distribuidora da Cartoon Network), em específico, não foram encontradas referências públicas online e, por questões de direito de uso e divulgação, não podem ser apresentadas aqui.

Em contrapartida, foram encontradas diretrizes semelhantes através da Netflix (2024) e *Prime Video* (2024), quanto ao uso das *Forced Narratives*. Tanto a Netflix quanto a *Prime Video* estabelece diretrizes que pautam pela exclusão de informações redundantes.

Tais diretrizes, no entanto, são claras para o uso das FNs, já que, em geral, são entregues em um arquivo à parte. Por outro lado, não fica totalmente claro que estratégia usar para o arquivo de legenda principal, que, a priori, e por via de regra, se legenda na íntegra, movendo-se o texto da legenda para a parte de cima (nesse caso), para não haver sobreposição com o texto em tela do letreiro.

Infere-se que a legendagem intralinguística de vídeos, como o exemplo acima da Cartoon Network, é pautada pela distribuidora. Ainda assim, cada distribuidor possui seu guia de estilo, com diretrizes, a fim de manter uma coerência no processo de produção, para fornecer a melhor experiência para o público.

Exemplos semelhantes também são encontrados não apenas no

*streaming*, mas também nas redes sociais. É o que se pretende discorrer brevemente na seção seguinte, na qual são apresentados alguns exemplos de vídeos com texto em tela narrado e legendas intralinguísticas.

#### **2.2.4 O Uso de Legendas e textos em telas nas Redes Sociais**

Enquanto as produções disponibilizadas pelas grandes distribuidoras são restritas aos produtores audiovisuais, conteúdos voltados às redes sociais vêm se popularizando em meio à produção e ao consumo de material digital. E, com o surgimento cada vez mais frequente de novas tecnologias, os usuários não apenas consomem, como também produzem e publicam seu conteúdo online (Scott, 2022), incluindo vídeos legendados, cujos parâmetros de legendagem são pautados pelo público em geral, ou seja, criadores de conteúdos independentes. No entanto, estes podem não seguir as mesmas regras de produção e não terem uma padronização quanto ao estilo de texto e formato de legenda inseridos. Em outras palavras, como traz Sasamoto (2024), cada criador de conteúdo tem a liberdade de escolher como inserir suas próprias legendas e textos em tela nos vídeos que produzem. A autora complementa, citando Johnson (2013), que, independentemente da forma como interagimos com o conteúdo, estamos em um momento de intensa representação textual, mas que, quanto mais texto é apresentado de forma simultânea na tela, maior o nível de processamento cognitivo exigido.

Sendo assim, uma vez que já foram apresentados definições e exemplos no segmento audiovisual no *streaming*, como podemos entender o uso dessa definição no segmento audiovisual das redes sociais? Para entender a justificativa desta pesquisa, faz-se necessário investigar exemplos reais de vídeos exibidos em plataformas de redes sociais, no intuito de identificar como eles são apresentados. Para tanto, leva-se em consideração aqui materiais inter ou intralinguísticos com texto em tela, desde que os vídeos possuam legendas, textos em tela e áudio no mesmo idioma, independentemente de ter havido um processo de tradução/localização em si, uma vez que o objeto de estudo é o vídeo final, e não o processo tradutório pelo qual o vídeo foi produzido.

Sendo assim, foram compilados alguns exemplos de vídeos disponíveis nas redes sociais, em especial, Instagram e YouTube. Parte-se da premissa de

que tais vídeos possuem legendas dinâmicas, e muitas delas assumem o caráter de texto em tela. Retomando as diretrizes vistas em Netflix (2024) e *Prime Video* (2024), acerca da exclusão de informações redundantes, as legendas desses vídeos são classificadas da seguinte forma:

- **Legendas Totais:** quando há a ocorrência tanto do texto da legenda quanto do texto em tela de forma simultânea;
- **Legendas Parciais:** quando o texto da legenda é excluído em momentos que há a presença de texto em tela cuja informação textual é a mesma;
- **Outros exemplos:** quando o padrão da legenda não se aplica aos estilos acima propostos.

Abaixo, são apresentados alguns materiais que ilustram as principais situações discutidas no escopo deste trabalho.

#### 2.2.4.1 Exemplos de Legendas Totais

A imagem abaixo apresenta dois exemplos do que é considerado como uso de legendas totais.

**Figura 11** – Captura de tela: Exemplo de Legenda Total



Fonte: (A) teachermatias, (B) smalladvantages.

Se tomarmos como exemplo a imagem da esquerda, tanto legenda quanto texto em tela apresentam a mesma informação: “Olha isso aqui”. O mesmo

acontece no exemplo B, à direita. Nesse caso, acredita-se, com base no princípio da redundância (Diao; Sweller, 2007, apud Liao, et al., 2021), que a alocação da atenção em duas fontes de informação redundantes pode diminuir a capacidade de entendimento do conteúdo, à medida em que a redundância de informações impactaria os padrões de leitura das legendas. Kruger, Hefer e Matthew (2014) complementam, tendo em vista que a competição entre fontes de informação também pode ter um impacto negativo e aumentar o esforço cognitivo. Vale lembrar que essa competição entre informações pode acontecer não apenas entre texto falado e escrito, mas também entre texto falado e imagens.

#### *2.2.4.2. Exemplos de Legendas Parciais*

Além dos exemplos acima, com o uso de legendas totais, também foram encontrados exemplos de legendas parciais, onde há a exclusão da legenda no momento em que a informação redundante do vídeo é apresentada. A imagem abaixo, retirada do canal *1 Minute Show*, apresenta uma sequência de frames exemplificando o uso de legendas em diálogo com textos em tela. Sua informação verbo-visual é a mesma informação narrada e, portanto, seria inserida como um arquivo de legenda em uma situação convencional. No entanto, os criadores do vídeo optaram por excluir do arquivo de legendas as informações textuais semelhantes às informações textuais inseridas através do texto em tela em formato de caixa de diálogo, usado possivelmente como elemento integrativo ao texto do vídeo (mensagem de texto). No exemplo, a caixa de diálogo permanece em tela ainda por alguns segundos mesmo após a próxima legenda surgir.

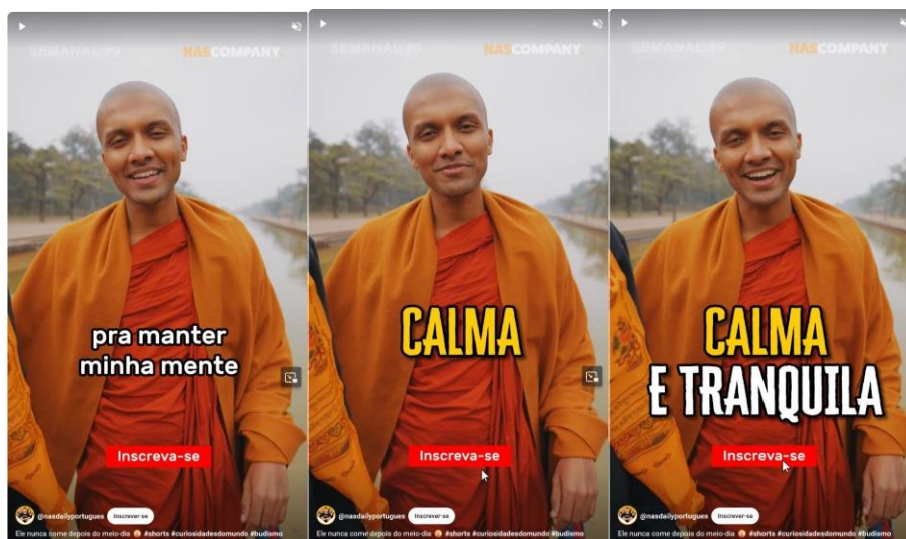
**Figura 12** – Captura de Tela: Sequência de imagens como exemplo de Legenda Parcial 1



Fonte: 1 Minute Show (2022).

No segundo exemplo, abaixo, retirado do canal Nas Daily Português (Nas Daily, 2024), o estilo de texto em tela pelas legendas dinâmicas muda um pouco em comparação ao primeiro exemplo. São apresentadas em fonte ampliada, no centro da tela, com uso de cores que dialogam com o design do vídeo. As legendas também seguem um padrão de cor semelhante. Além disso, também é usada a estratégia de legenda cumulativa, apresentada por Díaz Cintas e Remael (2021) na seção 2.1.1 deste trabalho.

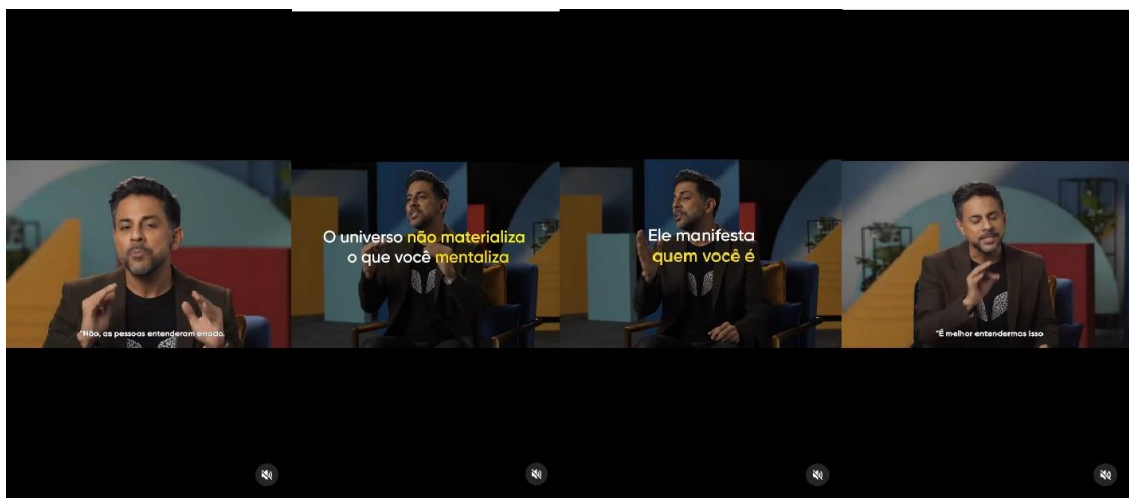
**Figura 13** – Captura de Tela: Sequência de imagens como exemplo de Legenda Parcial 2



Fonte: Nas Daily (2024).

Outro exemplo pode ser encontrado na página da Mindvalley Brasil, no Instagram. O vídeo em questão também segue o estilo de texto em tela na parte central, com cores que dialogam com o vídeo, e legendas convencionais na parte inferior da tela. Esse texto em tela aparece de forma gradual, palavra por palavra, diferentemente dos exemplos apresentados acima, com *Nas Daily* e *1 Minute Show*.

**Figura 14** – Captura de Tela: Sequência de imagens como exemplo de Legenda Parcial 3



Fonte: Mindvalley Brasil (2020).

Embora cada criador de conteúdo apresentado opte por inserir seu próprio estilo de legenda e animação, em todos os exemplos acima, é possível notar que, quando há o uso do texto em tela, como uma legenda dinâmica, a legenda principal é excluída, retirando-se, assim, a presença da informação redundante.

#### 2.2.4.3 Outros exemplos

Não obstante, nem sempre os dois casos apresentados acima – legendas parciais e legendas totais – se fazem presentes ou possíveis. Muitas vezes, tanto o uso das legendas como das informações textuais em tela como recurso de destaque são feitos de forma tão dinâmica, que se torna basicamente impossível adotar um padrão ou categorizá-lo em apenas duas definições – legendas totais ou legendas parciais. Em outras ocasiões, ambas as estratégias são aplicadas no mesmo vídeo. As imagens abaixo mostram exemplos retirados da página *Orube Game Studio*, no Instagram.

**Figura 15** – Captura de Tela: Sequência de imagens como exemplo de Legenda Parcial e Total

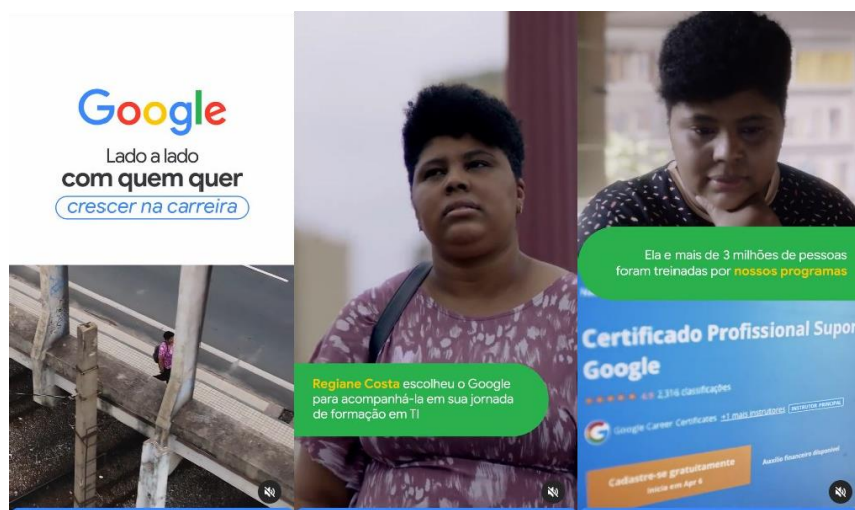


Fonte: Orube Game Studio (2023).

Logo na abertura no vídeo, a *Instagrammer* utiliza um texto em tela substituindo o uso da legenda, já que ambas as informações são as mesmas: “A sequência do Super Mombo Quest tá saindo”. Sendo assim, aqui, podemos atribuir o uso de legenda parcial. Na sequência, ela faz o uso de legendas convencionais verbatim e, no último frame da sequência, mantém a legenda quando utiliza o logotipo do jogo sobre o qual fala no vídeo. Nesse caso, pode-se dizer que houve o uso de legenda total, tendo em vista que há a repetição de informações semelhantes.

O próximo exemplo, retirado da página do Instagram da Google Brasil, foge um pouco aos estilos previamente apresentados aqui, uma vez que o vídeo não faz uso de legendas convencionais para transcrever o que é falado no áudio; pelo contrário, o faz totalmente através de textos em tela, com blocos de informação que surgem no estilo deslizante. Ao longo de todo o vídeo, as informações sonoras são transcritas com o uso de animações de vídeo para a inserção de bloco de texto, que mudam de cor, estilo, tamanho e posição na tela.

**Figura 16** – Captura de Tela: Sequência de imagens como exemplo de legenda em texto em tela dinâmico



Fonte: Google Brasil (2023).

Em suma, como pode ser observado a partir dos exemplos acima, nota-se uma tentativa, por parte de alguns criadores de conteúdo, de evitar o uso de informações redundantes em legendas, quando há a presença de textos em tela com informações verbo-visuais semelhantes. No entanto, outros exemplos trazem para a pauta que esse estilo nem sempre é aplicado, e ambas as informações são mantidas em tela de forma simultânea. Por outro lado, em decorrência da própria natureza dinâmica e interativa (Sasamoto, 2024; Johson, 2013) a que se objetivam as redes sociais, como Instagram e YouTube, vídeos com legendas produzidas e aplicadas de forma dinâmica também se fazem bastante presentes, não possibilitando uma categorização feita de forma objetiva.

Além disso, questiona-se como o uso excessivo de informação verbo-visual sem diretrizes para esse gênero audiovisual pode impactar o processo de recepção do público, o qual pode ser analisado através de metodologias que utilizem tecnologia de rastreamento ocular. Embora pesquisas nessa vertente ainda sejam inexploradas, a seguir, são discutidos alguns trabalhos com foco em rastreamento em legendagem e pesquisas exploratórias voltadas ao uso de texto em tela em outros gêneros audiovisuais.

## 2.3 Rastreamento ocular e legendagem

A leitura de legendas em contextos audiovisuais é um fenômeno que envolve interações complexas entre estímulos visuais, textuais e auditivos. Modelos de leitura em legendagem, como o proposto por d'Ydewalle e De Bruycker (2007), descrevem a leitura de legendas como um processo quase automático, em que os olhos dos espectadores tendem a se fixar imediatamente na área da legenda quando esta aparece na tela, independentemente de seu nível de compreensão do idioma. Essa leitura automática ocorre devido à posição fixa das legendas e à sua sincronia com os elementos visuais e auditivos do vídeo. O modelo também sugere que os espectadores frequentemente leem toda a legenda antes de explorar outros elementos visuais da tela, o que pode ser um indicativo de que legendas têm prioridade cognitiva.

Adicionalmente, Kruger, Heffer e Matthew (2014) propõem que o processamento de legendas é influenciado por fatores como tempo de exposição, sincronia com o áudio e complexidade textual. Segundo o modelo, legendas mal projetadas, como aquelas que incluem informações redundantes ou estão fora de sincronia com os elementos audiovisuais, podem aumentar o esforço visual e prejudicar a fluidez da leitura. Este estudo adota como base essas proposições teóricas para analisar de que forma diferentes designs de legendas, em vídeos com textos em tela, impactam os padrões de fixação e a compreensão do espectador.

Muitos estudos voltados à investigação dos processos cognitivos usando rastreamento ocular surgiram a partir dos anos 1980, embora tais estudos possam ser vistos como integrantes da terceira era de pesquisas com rastreamento ocular, datada desde meados dos anos 1970 (Rayner, 1998). Vale ressaltar que tais estudos tinham como foco a investigação de leitura de textos impressos ou estáticos (Kruger et al, 2020). Até então, não eram comuns pesquisas com foco na área audiovisual, como complementam Díaz Cintas e Remael (2021) e Araujo (2022), com abordagem em rastreamento ocular, principalmente, pelo fato de as legendas se caracterizarem como textos dinâmicos, competirem com as imagens integrantes do vídeo e competirem por recursos cognitivos com sons verbais e não verbais, como aponta Kruger *et al.*, (2020). E, embora os mesmos princípios relacionados às áreas de interesse para

objetos e imagens estáticos possam ser aplicados a vídeos, a investigação em materiais com vídeo requer a definição de áreas de interesse dinâmicas, já que esta é a natureza do material de análise, onde os elementos que o compõem também estão em movimento (Godfroid, 2019), como no caso de legendas.

Segundo Jensema et al (2000), os resultados de sua pesquisa demonstraram que adição de legendas gera uma mudança no movimento dos olhos. Em outras palavras, o espectador é levado a direcionar o olhar para a área do vídeo onde a legenda aparece, fazendo com que a leitura de legendas seja quase automática e obrigatória. Nesse processo, a pessoa não tem muito controle sobre sua atenção no que diz respeito a ler ou não um texto em tela. Esse achado também se alinha ao que sugerem d'Ydewalle e De Bruycker (2007).

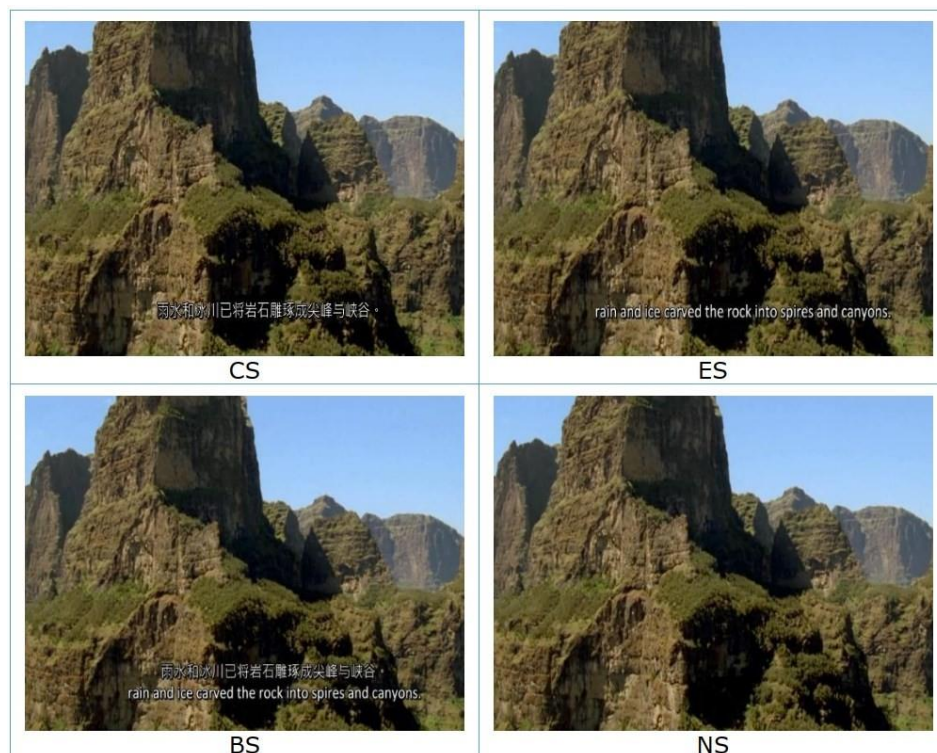
Dessa forma, é dedutível, então, que textos dinâmicos, como os exemplos apresentados acima, levam o público automaticamente a olhar para a região do vídeo em que o texto aparece. Assim, o uso ou não das legendas nesses casos – quando o texto em tela é o mesmo que o texto falado no áudio – geraria maior ou menor esforço cognitivo por parte de quem o assiste, já que se tratam de dois estímulos simultâneos para a atenção do espectador? Não foram encontrados trabalhos substanciais com foco específico na investigação do movimento ocular durante a leitura de legendas com texto em tela em contexto de redes sociais. No entanto, foram encontrados trabalhos voltados à área de legendagem em geral, e faz-se aqui, então, um recorte dessas pesquisas na busca de encontrar uma linha teórico-metodológica que possa ser usada como base para a presente pesquisa.

Como exemplo, temos a pesquisa desenvolvida por Liao, Kruger e Doherty (2020), cujo objetivo foi investigar o impacto de legendas bilíngues na distribuição da atenção, no esforço cognitivo e na compreensão. Os autores afirmam que, embora haja estudos com foco na investigação do impacto de legendas no esforço cognitivo, a fim de se entender o efeito que causam no processo de compreensão e aprendizagem, ainda pouco se sabe sobre o impacto de legendas bilíngues, já que a maioria dos estudos possui foco em legendas monolíngüísticas. O estudo se torna relevante para a presente pesquisa na medida em que busca analisar o esforço cognitivo mediante a presença de duas legendas simultaneamente.

Para o estudo, Liao, Kruger e Doherty (2020) usaram cinco vídeos de 5 minutos de duração cada, e participaram 17 falantes nativos de chinês como L1

(língua 1) e falantes de inglês como L2 (língua 2) (17 válidos, dos 20 participantes). O experimento foi conduzido em quatro condições experimentais para cada vídeo, sendo estas: (1) narrativa do vídeo em inglês com legendas em chinês; (2) narrativa do vídeo em inglês com legendas em inglês; (3) narrativa do vídeo em inglês com legendas tanto em inglês quanto em chinês; (4) narrativa do vídeo em inglês sem legendas. Cada participante foi submetido às quatro situações experimentais, conforme ilustrado na imagem abaixo.

**Figura 17** – Condições experimentais para cada vídeo da pesquisa de Liao, Kruger e Doherty (2020)



Fonte: Liao, Kruger e Doherty (2020, p. 75).

Além disso, os participantes também responderam a um questionário visando avaliar o esforço cognitivo, bem como foram solicitados a escrever o máximo de informações que conseguissem recuperar dos vídeos assistidos.

Para a análise dos dados da movimentação ocular, os autores usaram como medidas a porcentagem do tempo de permanência na AI, através do *Dwell Time %*. As AIs foram definidas em três regiões da tela, de acordo com cada elemento visual de análise. A Fig. 18, abaixo, ilustra a delimitação das AIs nas três regiões, como segue:

- Uma AI para a área total do vídeo;
- Uma AI para a área da legenda em inglês;
- Uma AI para a área da legenda em chinês.

**Figura 18** – Áreas de Interesse definidas por Liao, Kruger e Doherty (2020), em três regiões do vídeo



Fonte: Liao, Kruger e Doherty (2020, p. 78-79).

Os resultados demonstraram que os participantes pareceram passar tempo similar olhando para as legendas tanto em L1 quanto em L2, durante a exibição dos vídeos da condição bilíngue. No entanto, os autores ressaltam que os participantes têm a tendência de escolher uma das legendas (L1 ou L2) como fonte de informação verbo-visual dominante durante o processo de leitura. Além disso, os autores inferem que os participantes têm a tendência de ajustar o padrão de visualização para focar na principal fonte de informação.

Em outro estudo realizado por Liao *et al.* (2021), intitulado “*Using Eye Movements to Study the Reading of Subtitles in Video*” (Usando Movimentos Oculares para Estudar a Leitura de Legendas em Vídeo), os autores, conforme indicado pelo título do trabalho, buscaram investigar as consequências geradas a partir da leitura de legendas queimadas em vídeo, a partir da presença de duas variáveis de pesquisa, a saber: se a presença ou ausência de informação imagética de fundo e a velocidade de leitura das legendas influenciam na compreensão textual. Como objetivo secundário, os autores buscaram avaliar as estratégias adotadas pelos participantes da pesquisa para compreender o texto apresentado. Para investigar tais objetivos, os autores utilizaram duas condições em vídeos:

1. vídeos com a presença de informações visuais de fundo;
2. vídeos sem a presença de informações visuais de fundo;

3. vídeos com legendas nas velocidades de 12, 20 e 28 CPS - caracteres por segundo.

**Figura 19** – Exemplo de estímulo usado por Liao et al. (2021), com a presença (imagem à direita) e a ausência (imagem à esquerda) de informação visual de fundo



Fonte: Liao et al. (2021).

No total, participaram da pesquisa 31 falantes de inglês como idioma nativo, que assistiram às seis condições – sendo, no entanto, em vídeos diferentes, ou seja, um vídeo para cada variável – e responderam a um questionário sobre os vídeos a fim de avaliar o nível de compreensão sobre cada vídeo apresentado.

Para avaliar as condições presentes, os autores optaram por utilizar três variáveis universais, com áreas de interesse definidas na totalidade da região na qual as legendas estavam presentes:

- (1) duração média das fixações;
- (2) duração média das sacadas;
- (3) número total de fixações.

Para investigar a mudança de atenção entre as legendas, também foram utilizados como parâmetro de análise:

- (4) número de *crossovers*;
- (5) *gaze duration*;
- (6) tempo total;
- (7) probabilidade de pular uma palavra.

Os resultados de Liao et al. (2021) indicaram que a compressão dos vídeos foi maior entre os participantes quando havia a presença das informações visuais de fundo, opondo-se à hipótese de que várias fontes de informação podem gerar menor compreensão. No que diz respeito à duração média de fixações, as

durações foram mais longas na região do vídeo – quando houve a presença de informação visual de fundo –, indicando que os participantes demonstraram atenção ao conteúdo do vídeo. Ao se analisar a duração média de sacadas, obteve-se que os participantes apresentaram sacadas mais longas na presença do vídeo. Já no que diz respeito ao número total de fixações, ocorreram mais fixações nas legendas do que na região do vídeo, mesmo na presença de fonte de informação visual de fundo, reduzido com o aumento da velocidade de leitura das legendas. O número de *crossovers* entre legendas e vídeo aumentou com a presença de informações visuais de fundo e diminuiu à medida que a velocidade das legendas aumentou.

Em resumo, os dados obtidos por Liao *et al.* (2021) indicam que os participantes são suscetíveis a mudanças presentes no conteúdo, como foi o caso dos vídeos em questão. Ainda segundo os autores (p. 430),

Por exemplo, as nossas análises globais mostram que, na presença de vídeo simultâneo, os leitores passam mais tempo examinando esse conteúdo. Este comportamento provavelmente reflete o interesse dos participantes no vídeo e talvez o uso estratégico de seu conteúdo para complementar a compreensão das legendas (Liao *et al.*, 2021, p.430, tradução nossa)<sup>7</sup>.

Como pode ser observado a partir das pesquisas apresentadas acima, há uma tendência à investigação e análise de legendas intralinguísticas, análise comparativa entre L1 e L2, cada qual com seu foco investigativo. Diferentemente de pesquisas feitas com textos estáticos, cuja velocidade de leitura é pautada pelo próprio leitor, em se tratando de materiais em vídeo, cujo texto está em constante movimento, o ritmo de leitura será definido pela velocidade de exibição das legendas (Kruger *et al.*, 2020). Sabe-se também que o processamento de legendas é afetado pela velocidade em que as legendas são apresentadas, além de outros fatores, como número de linhas, formatos de apresentação das legendas e sua legibilidade, presença ou não do áudio no vídeo (Szarkowska; Gerber-Morón, 2018), bem como fatores pertinentes ao próprio sujeito, como nível de leitura (Sasamoto, 2024; Orero; 2018).

Por serem fatores diretamente relacionados ao processamento das legendas, muitas pesquisas são voltadas à investigação de parâmetros como velocidade de leitura e segmentação linguística (Sasamoto 2024), seja por parte

de espectadores surdos ou ouvintes ou ambos, a fim de obter parâmetros que auxiliem no entendimento dos padrões de leituras. Exemplos desses trabalhos são Assis (2021), Szarkowska e Gerber-Morón (2018), Carvalho e Seoane (2019), Vieira (2016), Kruger, Wisniewska e Liao (2022).

Há também trabalhos com foco investigativo na leitura de legendas destinadas a contextos acadêmicos/educacionais e aquisição de L2, como Kruger (2013) e Matthew (2020); e, dentre os fatores também apresentados acima, resumidos por Sasamoto (2024), há pesquisas voltadas à investigação da influência da presença ou não de áudio no processamento de legendas, como Liao *et al.* (2022) e Lâng (2016).

Além dos fatores técnicos, outro aspecto em pesquisas com rastreador ocular é a legibilidade, que inclui taxas de leitura, hábitos de leitura, complexidade do texto, carga semântica, mudanças de cena e velocidade de fala (Gambier, 2003, p. 56 *apud* Sasamoto, 2024, p. 106-107, tradução nossa<sup>8</sup>).

Em sua maioria, essas pesquisas têm como objeto de estudo legendas, cujos parâmetros são bem definidos – ou seguem o modelo tradicional do que conhecemos como legendas – e que fornecem dados importantes sobre o processamento de legendas. Contudo, de que forma essas metodologias podem ser aplicadas a trabalhos de cunho investigativo com foco em vídeos com legendas de natureza criativa/dinâmica não-convencional, como os exemplos apresentados acima ou as legendas de impacto, apresentadas por Sasamoto (2024) e O'Hagan e Sasamoto (2016)?

O'Hagan e Sasamoto (2016) buscaram investigar o uso de “legendas de impacto” no contexto da tradução audiovisual como prática inovadora, argumentando que esse estilo de legenda é pouco conhecido fora da Ásia, o que corrobora para o que já foi mencionado sobre a escassez de trabalhos dedicados. As autoras complementam que a maioria dos estudos são de base teórica, com poucos estudos de recepção. O estudo utilizou um trecho de aproximadamente 22 minutos de um programa de TV japonês, e participaram 12 estudantes universitários japoneses

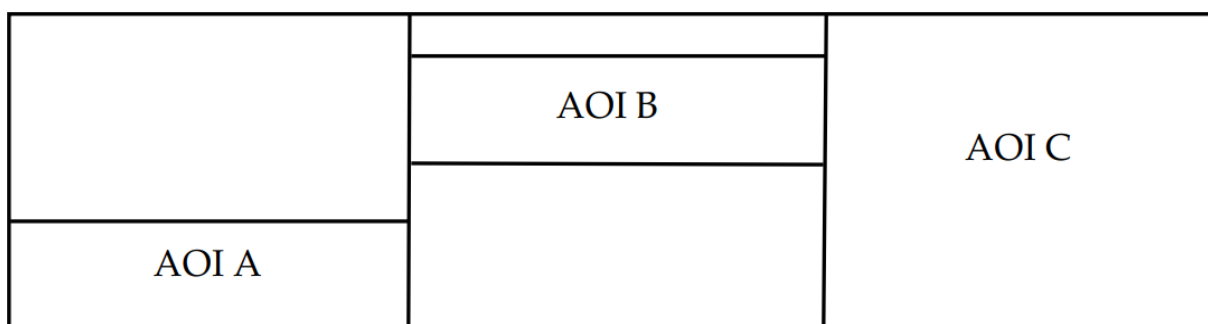
---

<sup>7</sup>The eye-movement record provides tentative answers to these questions. For example, our global analyses show that, with concurrent video content, readers spend more time examining that content. This behavior likely reflects participants' interest in the video, and perhaps the strategic use of its content to supplement subtitle comprehension.

<sup>8</sup>Another aspect of eye-tracking research is readability, which includes reading rates, reading habits, text complexity, semantic load, shot changes, and speech rates (Gambier, 2003, p. 56 *apud* Sasamoto, 2024, p. 106-107).

nativos (com dados válidos). Para a apresentação dos estímulos, foi usado um monitor de TV de 32 polegadas, e o rastreador ocular *Tobii Glasses* foi escolhido para o monitoramento dos movimentos oculares. Além disso, três áreas de interesse foram estabelecidas, levando em conta o espaço do monitor e da natureza dinâmica das legendas, ou seja, tempo, posição e tamanho variados, de acordo com a metodologia adaptada de Josephson e Holmes (2006).

**Figura 20** – Áreas de Interesse definidas por O'Hagan e Sasamoto (2016)



Fonte: O'Hagan e Sasamoto (2016, p. 14).

Com base nas AIs definidas, as autoras traçaram as seguintes hipóteses:

(1) os olhos frequentemente se fixam na região dos rostos na parte superior central da tela; (2) as visitas à região inferior, onde legendas de impacto são colocadas, não são frequentes (O'hagan; Sasamoto, 2016, p. 14).

Os resultados indicaram que participantes familiarizados com essas legendas fizeram a leitura delas, mas não se fixaram nelas e concentrando-se nas regiões centrais da tela onde os rostos aparecem. Ainda segundo as autoras, embora os resultados não tenham levado a conclusões definitivas sobre o impacto das legendas nos espectadores, o estudo piloto indicou direções futuras para experimentos controlados, e que o uso de rastreamento ocular demonstra potencial para informar produtores sobre a recepção fisiológica dos espectadores às legendas, influenciando estratégias futuras.

Josephson e Holmes (2006) também buscaram investigar se espectadores acostumados com a presença de textos em tela em noticiários conseguiriam lidar com o excesso de informações. Tal pesquisa se torna relevante para o presente trabalho, na medida em que fornece informações experimentais e teóricas de base. No estudo, foram formuladas quatro perguntas de pesquisa com foco em

atenção visual e capacidade de memorização.

RQ1: A distribuição da atenção visual para áreas da tela durante a visualização de notícias na TV está associada à presença de textos em tela realçadores? RQ2: As similaridades da trajetória ocular estão associadas à presença de textos em tela realçadores? RQ3: A quantidade de informações lembradas no conteúdo sonoro dos noticiários na TV está associada à presença de textos em tela realçadores? RQ4: Os espectadores preferem design de noticiários na TV com textos em tela realçadores? (Josephson; Holmes, 2006, p. 156, tradução nossa<sup>9</sup>).

Participaram da pesquisa 36 universitários, que assistiram a três noticiários gravados de aproximadamente 2 minutos de duração. Cada noticiário teve três versões manipuladas acrescentadas de informações visuais e textuais, uma com o noticiário convencional, uma com a presença de “legendas deslizantes (*crawler*)” e outra versão com a manchete da matéria e legendas deslizantes. Após assistirem aos vídeos, os participantes responderam a um questionário pós-coleta com perguntas sobre os vídeos e foram perguntados acerca de qual versão lhes agradara mais. Os autores analisaram os dados com base na distribuição da atenção, na comparação da trajetória ocular entre os participantes e nas respostas fornecidas no questionário pós-coleta. Os resultados indicaram que a adição dos textos em tela levou a uma mudança quanto a que parte da tela os sujeitos olharam sem prejudicar, entretanto, a compreensão do conteúdo. Por fim, ao responder às perguntas de pesquisa, as respostas dos participantes mostraram que, “embora os três designs tenham sido avaliados favoravelmente, há dois campos de preferência: os que preferem o design padrão aos designs aprimorados, e os que preferem os designs aprimorados ao design padrão (Josephson; Holmes, 2006, p. 161, tradução nossa<sup>10</sup>)”.

Neste capítulo, foram apresentadas algumas discussões acerca do uso de rastreamento ocular em pesquisas com vídeos legendados em algumas fontes, a saber: pesquisas convencionais usando vídeos tradicionais com foco em variáveis – como número de linhas, segmentação ou presença simultânea de legendas em dois idiomas – e pesquisas com foco em vídeos com predominância

---

<sup>9</sup> RQ1: Is distribution of visual attention to screen areas during TV news viewing associated with the presence of on-screen enhancements? RQ2: Are eye-path similarities associated with the presence of on-screen enhancements? RQ3: Is amount of recall of information in the audio content of TV news stories associated with the presence of on-screen enhancements? RQ4: Do viewers prefer TV news designs with on-screen enhancements?

<sup>10</sup> While the three designs were favorably evaluated, there are two preference camps: those who like the standard design more than enhanced designs, and those who like enhanced designs more than the standard design.

de textos em tela – telop, legendas de impacto –, cujas legendas/elementos visuais fogem do padrão comum do que se conhece como “legendagem”. Em suma, essas pesquisas objetivam estudar aspectos controláveis na produção de legendas que podem influenciar na sua recepção, como: velocidade de legendas, número de linhas, segmentação, mudança de cena, frequência das palavras, uso ou não de canal sonoro etc. No entanto, nota-se ainda a necessidade de se explorar melhor outros fatores controláveis em vídeos que são populares, mas não possuem pesquisas dedicadas ao uso não regulamentado de texto em tela, e se tal fato pode afetar a experiência do público.

Sendo assim, apresentam-se, a seguir, algumas das principais medidas de investigação, compiladas a partir das leituras feitas e adaptadas para o presente trabalho.

## 2.4 Hipóteses

A partir dos trabalhos de Liao, Kruger e Doherty (2020) e Liao et al. (2021), descritos acima, cujos focos de investigação são pautados na leitura de legendas, entende-se a importância do tempo de permanência da área de interesse, que pode ser obtido a partir do *Dwell Time*.

Esta análise pode ajudar a fornecer indícios sobre diversos pontos de análise, como resume Holmqvist *et al.* (2011), a saber:

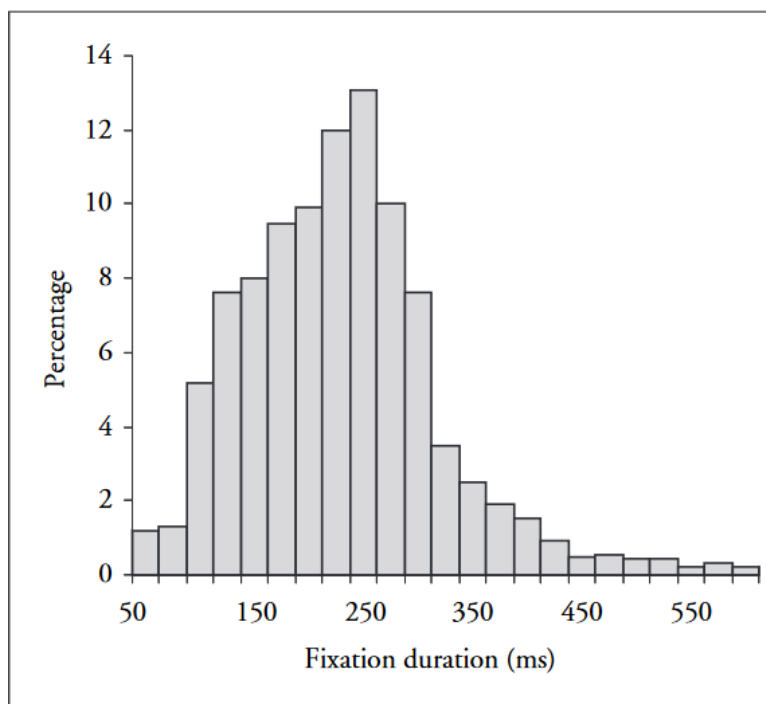
- interesse e informatividade;
- incerteza e pouca consciência situacional;
- dificuldade em extrair informação geral e ao nível de palavra;
- escolha/preferência consciente de determinado elemento.

Sendo assim, um maior tempo de permanência na área de interesse pode indicar as situações acima listadas. De acordo com O'Hagan e Sasamoto (2016), conforme mencionado, os espectadores têm a tendência de olhar para a região central da tela. Dessa forma, considerando-se a presença de uma legenda e um texto em tela, espera-se que o tempo de permanência seja relativamente distribuído entre a legenda e o texto em tela. Do contrário, quando há apenas o texto em tela, espera-se que haja tempo de permanência maior nessa região, dada a ausência de legendas na região inferior ao vídeo.

Outra medida comumente utilizada é a duração média de fixações. Cada

fixação pode ser entendida como um ponto em que o movimento do olho se encontra relativamente estável sobre dado elemento, seja este uma palavra ou outro elemento visual (Assis, 2021; Holmqvist *et al.*, 2011). O valor médio de duração de fixações por sujeitos adultos varia entre 200 ms e 250 ms (Rayner; Juhasz; Pollatsek, 2005; (Rayner, 1998), como ilustra a Fig. 21, abaixo, podendo haver fixações mais curtas, em torno de 50 ms, bem como fixações mais longas, acima de 550 ms, onde fixações mais curtas, de até 100 ms, podem indicar varredura rápida por parte do leitor na busca de informações e geralmente está associada à sacada.

**Figura 21 – Distribuição de duração de fixações**



Fonte: Rayner (1998), adaptado por Rayner; Juhasz e Pollatsek (2005, p. 81).

Ainda segundo Rayner (1998), essas durações médias variam de acordo com a atividade de leitura com a qual se está envolvido. O Quadro 3, abaixo, traduzido de Rayner (1998), resume algumas dessas atividades e suas respectivas durações.

**Quadro 3** – Duração média das fixações em leitura, busca visual e percepção de cena (adaptado)

| Tarefa                   | Duração Média das Fixações (ms) |
|--------------------------|---------------------------------|
| Leitura silenciosa       | 225                             |
| Leitura em voz alta      | 275                             |
| Busca visual             | 275                             |
| Percepção de cena        | 330                             |
| Leitura de letra musical | 375                             |
| Digitação                | 400                             |

Fonte: traduzido e adaptado de Rayner (1998, p. 373).

Com base nos valores acima apresentados e de acordo com Assis (2021, p. 38), “a leitura silenciosa, com média de duração de 225 milissegundos, seria a mais próxima ao contexto da leitura dinâmica da legendagem”.

Em conjunto com a duração média das fixações, tem-se o número total de fixações: o número de fixações pode indicar uma relação direta com a média de fixações, já que são variáveis dependentes, ou seja, um maior número de fixações pode indicar mais engajamento durante a leitura e reduzir o tempo médio de fixações, indicando uma leitura mais fluída e com menos ruídos.

Com base nessas duas médias, é esperado que, em contextos de legendas totais, como exemplificado na seção 2.2.4.1, a duração média das fixações e o número de fixações sejam relativamente distribuídos entre as duas AIs; em contexto de legendas parciais, espera-se que a duração média de fixações e o número de fixações sejam maiores na AI 2. De modo geral, espera-se que a duração média de fixações na AI 2 na condição 2 sejam menores do que os valores obtidos na AI 1 na condição 1, bem como haja mais fixações na AI 2 na condição 2 em comparação à AI na condição 1.

Cita aqui também, o tempo da Primeira Fixação na Área de Interesse: embora essa medida não seja considerada como inteiramente válida para a análise de processamento de leitura – uma vez que, na leitura de legendas, a primeira fixação costuma ser direcionada ao centro da legenda antes de uma regressão para o início da legenda (Kruger; SteYN, 2014) –, a utilização da primeira fixação se torna relevante aqui, na medida em que está aliada à alocação de atenção e se utiliza-se valores da primeira fixação na área de interesse.

Segundo Holmqvist *et al* (2011), em pesquisas envolvendo processamento durante a leitura, a primeira fixação em uma palavra pode estar associada com ativação lexical. Ainda segundo os autores, essa medida indica o processamento da primeira fonte de informação na qual o olho fixou, podendo haver maior ou menor grau de dificuldade de processamento por parte do sujeito, onde um tempo maior da primeira fixação pode indicar que a informação apresenta maior grau de complexidade de entendimento.

Em se tratando da primeira fixação na área de interesse, esta pode indicar o tempo gasto para processamento rápido, como reconhecimento e identificação de elementos visuais/textuais dentro da área de interesse fixada. Em suma, a duração da primeira fixação pode funcionar como uma medida na aquisição de informação visual da área fixada (Holmqvist *et al.*, 2011) e indicar o quão rapidamente o sujeito pode se engajar com a informação visual apresentada, o que pode ser útil para testar a eficácia do posicionamento e do formato das legendas.

Para essa medida, espera-se que, em vídeos com legendas totais, tanto na condição 1 quanto na condição 2, a duração média da primeira fixação seja mais longa na AI 2, tendo em vista o surgimento do texto em tela como fator desencadeante de alocação de atenção. Já em vídeos com legendas parciais, espera-se que a duração média da primeira fixação na AI 2 na condição 2 seja menor do que a duração média na AI 2 na condição 1, considerando que a condição 1 gere maior disputa de atenção, justificando uma média maior.

Adiciona-se também a análise de Sequência de Fixações: esta análise observa a ordem e o padrão das fixações entre as legendas, oferecendo *insights* sobre como o leitor alterna sua atenção e gerencia a leitura dentre as áreas de interesse. A Fig. 22, abaixo, seguida pelo quadro 4, traduzida do guia de utilização do Eyelink (2024, p. 141), ilustra um exemplo da análise de fixações na frase “*He said: “Truth is a pathless land”*”.



No Quadro 4, acima, cada coluna indica o número de fixações recebidas na área de interesse, originadas a partir de todas as possíveis áreas de interesse. Cada linha indica o número de fixações que se originaram a partir da área de interesse que a linha representa, e foram direcionadas a cada uma das demais áreas de interesse, indicada na linha (Eyelink, 2024, p. 141).

Ainda segundo o guia do Eyelink (2024), a fim de facilitar a compreensão, como exemplo, mostra-se a área de interesse AI 6 destacada, “pathless”. A AI recebeu ao todo 4 fixações, vindas da AI 3 (fixação 4), AI 4 (fixação 8), AI 7 (fixação 10) e da própria AI 6 (fixação 5); e dela partiram 3 fixações, fixação 5, (saída da própria AI 6), fixação 6 (AI 3) e a fixação 9 (AI 7). Ao se observar a quantidade de fixações recebidas, a AI 6, destacada em cinza na coluna, apesar de ter recebido 4 fixações, obteve somente uma delas originadas dentro da própria área de interesse, destacada com a cor cinza no cruzamento da linha com a coluna, ou seja, a fixação 5, feita a partir da sacada após a fixação 4, mostrada na Fig. 22.

Semelhantemente, ao se tomar como referência apenas duas áreas de interesse, AI 6 e AI 7, por exemplo, e tendo em mente que o presente trabalho analisa a sequência de fixações entre duas áreas de interesse, esta seria a situação:

**Quadro 5 – Matriz de fixações entre AI 6 e AI 7 adaptado de Eyelink (2024)**

| AI | IA FSA COUNT 6 | IA FSA COUNT 7 |
|----|----------------|----------------|
| 6  | 1              | 1              |
| 7  | 1              | 0              |

Fonte: adaptado de Eyelink (2024, p. 142).

No cruzamento de fixações, tem-se que duas fixações partiram da AI 6, ou seja, duas sacadas entre a área de interesse 6 e 7. Uma delas para a própria AI 6 e uma para a AI 7. Em outras palavras, não houve fixações originadas de sacadas dentro da própria AI 7, mas houve uma fixação após sacadas originadas dentro da própria AI 6. Em relação à AI 7, houve apenas uma fixação saída de AI 7 para AI 6.

Com base nisso, espera-se que, em vídeos com legendas totais, haja mais

transições entre as áreas de interesse AI 1 e AI 2, com menor número de fixações originadas dentro na própria AI; já em vídeos com legendas parciais, que haja menos transições entre as áreas de interesse AI 1 e AI 2, com maior número de fixações originados da própria AI.

Sendo assim, a análise da sequência de fixações pode ajudar a entender o número de fixações que foram seguidas por transcrições entre as áreas de interesse. Dentre as aplicações desta medida, Holmqvist *et al.* (2011) resume três aplicações para as quais o número de transições entre duas áreas de interesse tem sido usado como uma medida de avaliação, a saber: design, expertise e importância de determinada área. Ainda segundo os autores, quando uma área de interesse se apresenta como mais importante, a ponto de manter maior nível de monitoramento ou atenção, o número de transições é menor. Observa-se, segundo os autores acima, que o uso da análise da sequência de fixações pode ser válido para investigar um design mais adequado de legendas de acordo com as condições experimentais apresentadas, bem como obter indícios de qual área de interesse pode ser percebida como mais importante por parte dos sujeitos.

Dessa forma, a medida de transições entre as duas áreas de interesse (AI) pode ajudar a avaliar também a eficácia do design das legendas ou de qualquer elemento visual que as duas AIs representem. Pode-se concluir que um design eficaz deve facilitar uma leitura fluida e equilibrada, sem causar confusões ou sobrecargas desnecessárias de atenção em uma área. Se uma das áreas de interesse tem significativamente mais transições, isso pode sugerir que o design dessa AI está atraindo mais atenção, o que pode ser tanto um indicativo de sucesso (se for intencional) quanto um problema (se causar distração do foco principal).

A análise de sequência de fixações, portanto, não apenas contribui para compreender como os participantes interagem com as áreas de interesse, mas também pode fornecer dados iniciais para otimizar o design e a disposição de conteúdos legendados com textos em tela. Isso, com base em como as informações são processadas visualmente quando se tem uma legenda e um texto em tela no mesmo momento de exibição.

Com base nessas discussões, como investigar, então, a recepção de legendas de forma a contemplar as variáveis pretendidas, uma vez que diversos

fatores externos ao pesquisador estão em pauta? É o que se pretende discorrer nas próximas seções, onde se apresenta a proposta metodológica elaborada com base nas leituras feitas até o momento e nos trabalhos relacionados encontrados.

### 3 DESENHO EXPERIMENTAL

Este capítulo objetiva-se a apresentar a problemática identificada através das leituras mencionadas acima, a razão da proposta metodológica, e discutir sua pertinência para o tipo de pesquisa aqui abordada.

Segundo Orero *et al.* (2018), em decorrência de diversos fatores variáveis dos participantes e aspectos cognitivos, linguísticos e socioculturais, a criação de experimentos em TAV torna-se uma tarefa difícil, e manter o controle de todas as variáveis em grupos de controle nem sempre é possível. Uma alternativa, ainda segundo a autora, é criar grupos de controle com base em características específicas, que, embora possam comprometer a validade dos dados, podem ser úteis para a execução de pesquisas-piloto para o desenho experimental. O mesmo princípio pode ser aplicado à seleção dos vídeos, tendo em vista que – assim apresentado na seção de exemplos – esses materiais não possuem um guia padrão de estilo pelas razões previamente já discutidas.

Em resumo, a partir da análise feita nos materiais disponíveis, vimos que as legendas diferem em cor e tamanho da fonte, posição em que aparecem no vídeo, modo de exibição, como legendas deslizantes, blocos etc., bem como o uso de legendas totais ou parciais, conforme apresentado na seção anterior. Tais fatores dificultariam o processo de análise em decorrência do grande número de variáveis, além de variáveis não controláveis, caso se optasse por utilizar vídeos prontos retirados do YouTube ou Instagram, como a velocidade de legenda (CPS), o número de linhas, os elementos visuais nos vídeos, além do texto de legenda e textos em tela, tamanho das palavras etc.

Com base nisso, pensou-se na proposta de uma metodologia de coleta de dados voltada a vídeos em formato para redes sociais, cujo percurso metodológico é descrito a seguir e cuja linha de pesquisa enquadra-se nos Estudos da Tradução com foco em Tradução Audiovisual em pesquisa de Linguagem, Cognição e Recursos Tecnológicos.

### **3.1 Tipo de Pesquisa**

A descrição acima encerra tudo o que foi encontrado sobre legendagem e vídeos com texto em tela e casos que alguns autores chamam de redundantes. Contudo, até a presente data, não foram encontrados trabalhos voltados à investigação especificamente do movimento ocular durante a exibição de vídeo com texto em tela em contexto de redes sociais, ou seja, com uma legenda total e outra parcial. Diante do ineditismo, a presente pesquisa configura-se como um estudo piloto experimental de natureza descritiva explicativa. Trabalhos como os desenvolvidos por Liao, Kruger e Doherty (2020), O'Hagan e Sasamoto (2016) e Josephson e Holmes (2006) são usados como ponto de partida para a construção do design experimental aqui proposto.

### **3.2 Preparação do Material**

Inicialmente, pensou-se em utilizar vídeos já disponibilizados em redes sociais, como YouTube e Instagram, a partir da delimitação do uso de legendas nas duas categorias previamente discutidas. Isso poderia fornecer uma melhor validade ecológica dos dados, pois seriam retirados de contextos reais de uso. No entanto, como garantir que a velocidade das legendas, posição, cor, tamanho de fonte, bem como outras informações que competem pela atenção visual no vídeo não afetariam os resultados que se pretende analisar? Sendo assim, optou-se por construir os estímulos, desde a seleção dos temas a serem discutidos até a produção das legendas.

#### **3.2.1 Seleção dos vídeos**

Para o assunto dos vídeos, delimitaram-se três principais temas, por serem recorrentes da atualidade, sendo que, para cada tema, dois vídeos foram gravados, totalizando seis vídeos.

3.2.1.1 Tema 1: Trabalho e produtividade;

3.2.1.2 Tema 2: Transtorno do Déficit de Atenção e Hiperatividade (TDAH);

3.2.1.3 Tema 3: Saúde e Condicionamento Físico.

### 3.2.2 Controle de variáveis nos vídeos

Mesmo com os temas selecionados, como controlar os elementos textuais variáveis, como duração dos vídeos e número de palavras, tamanho das palavras e frequência em relação ao léxico do português? Para completar tal etapa, os roteiros que serviram para a gravação dos vídeos também foram manipulados e criados com a ajuda de inteligência artificial. Os roteiros precisavam ter um número médio de palavras equivalentes, para que se mantivesse também uma duração média aproximada.

Esse processo foi constituído de duas etapas principais:

3.2.2.1 Elaboração dos roteiros com ajuda de inteligência artificial;

3.2.2.2 Validação dos roteiros para as variáveis selecionadas.

Para a preparação de cada roteiro, foi estabelecido um limite máximo de 200 palavras, a fim de que cada vídeo não ultrapassasse 90 segundos de duração. Os roteiros continham entre 162 e 179 palavras, com número médio de 168 palavras. A tabela abaixo resume o número de palavras final de cada roteiro, após serem manipulados, sobre os temas pretendidos.

**Tabela 1 – Número de palavras dos roteiros**

|                          | Roteiro 1 | Roteiro 2 | Roteiro 3 | Roteiro 4 | Roteiro 5 | Roteiro 6 | Média |
|--------------------------|-----------|-----------|-----------|-----------|-----------|-----------|-------|
| Nº de Palavras           | 170       | 162       | 179       | 171       | 162       | 164       | 168   |
| Duração dos vídeos (seg) | 71        | 81        | 88        | 84        | 80        | 85        | 82    |

Fonte: elaborado pelo autor.

Após a criação dos roteiros, os textos foram exportados em formato .txt e processados no AntConc (Anthony, 2023). O AntConc é uma ferramenta de análise de *corpus*, através da qual foi possível obter uma listagem de todas as palavras de cada roteiro, ordenadas de acordo com a frequência em que aparecem no roteiro processado. A tabela abaixo apresenta as três palavras mais frequentes de cada roteiro. Os roteiros podem ser conferidos na seção dos apêndices.

**Tabela 2** – Frequência das principais palavras de cada roteiro

| <b>Roteiro</b> | <b>Entrada</b> | <b>Frequência no <i>Corpus</i></b> |
|----------------|----------------|------------------------------------|
| Roteiro 1      | sono           | 7                                  |
|                | dormir         | 5                                  |
|                | dica           | 3                                  |
| Roteiro 2      | trabalho       | 7                                  |
|                | você           | 4                                  |
|                | espaço         | 3                                  |
| Roteiro 3      | TDAH           | 7                                  |
|                | pessoas        | 4                                  |
|                | é              | 4                                  |
| Roteiro 4      | mentalidade    | 4                                  |
|                | é              | 4                                  |
|                | aquecimento    | 3                                  |
| Roteiro 5      | saúde          | 5                                  |
|                | número         | 4                                  |
|                | bem            | 3                                  |
| Roteiro 6      | dica           | 3                                  |
|                | saudável       | 3                                  |
|                | saúde          | 3                                  |

Fonte: elaborado pelo autor.

Após processadas no AntConc, as listas foram divididas entre palavras de conteúdo, como verbos, adjetivos e substantivos, e palavras de não-conteúdo, como preposições, artigos e conjunções. Cada palavra de conteúdo foi processada no Léxico da Língua Portuguesa (Lexporbr, 2019), a fim de se obter:

- a frequência das palavras em relação ao léxico do português, representada por “zipf\_escala”;
- o tamanho de cada palavra, representado por “nb\_letras”, com base no número de caracteres.

Os dados da análise podem ser vistos na tabela abaixo:

**Tabela 3 – Média da frequência e tamanho das palavras dos roteiros**

|                    | Roteiro 1 | Roteiro 2 | Roteiro 3 | Roteiro 4 | Roteiro 5 | Roteiro 6 | Média |
|--------------------|-----------|-----------|-----------|-----------|-----------|-----------|-------|
| <b>zipf_escala</b> | 4,37      | 4,37      | 4,43      | 4,28      | 4,38      | 4,24      | 4,34  |
| <b>nb_letras</b>   | 6,75      | 6,93      | 7,68      | 7,15      | 7,02      | 7,21      | 7,12  |

Fonte: elaborado pelo autor.

A partir da análise da tabela, pode-se perceber que a listagem de palavras apresenta uma construção homogênea, com palavras frequentes e de tamanho aproximado, o que pode levar a menores interferências durante o processo de leitura e processamento das legendas dos vídeos no experimento.

Após ter roteiros mediamente homogêneos e validados em frequência e tamanho de palavras, partiu-se para a gravação dos roteiros em vídeo, assim como descrito a seguir.

### **3.2.3 Gravação dos vídeos**

Cada vídeo foi gravado com um celular no formato vertical 9x16 (*video ratio*), por ser o formato mais comum nas redes sociais, como os *reels* do Instagram – modalidade de vídeos analisados anteriormente, na seção 2.2.4. Com intuito de manter maior homogeneidade entre os vídeos usados no experimento, optou-se por gravá-los sobre fundo neutro e vestimenta de cor preta, para melhor apresentação visual das legendas, bem como para ter o

mínimo de elementos visuais possíveis. Os roteiros foram processados no aplicativo *Elegant Teleprompt*, que permite importar os roteiros, ajustar a velocidade de exibição do texto e sobrepô-lo ao aplicativo da câmera do celular. Dessa forma, foi possível manter um melhor controle sobre a velocidade de leitura dos roteiros e, conseqüentemente, garantir uma velocidade de legenda mais homogênea entre os vídeos. Em suma, os vídeos ficaram com durações entre 71 e 88 segundos.

**Tabela 4 – Número de palavras dos roteiros**

|                                 | Roteiro 1 | Roteiro 2 | Roteiro 3 | Roteiro 4 | Roteiro 5 | Roteiro 6 | Média |
|---------------------------------|-----------|-----------|-----------|-----------|-----------|-----------|-------|
| <b>Duração dos vídeos (seg)</b> | 71        | 81        | 88        | 84        | 80        | 85        | 82    |

Fonte: elaborado pelo autor.

### 3.2.4 Preparação das Legendas

Uma vez gravados os vídeos, passou-se à etapa de legendagem. No capítulo teórico, observou-se que fatores como velocidade das legendas, número de linhas e segmentação podem influenciar na recepção do material. Sendo assim, aqui se buscou produzir as legendas de forma a excluir essas variáveis dos pontos de análise. A utilização do *Elegant Teleprompt* mostrou-se útil no controle da velocidade de leitura do texto, resultando em legendas cujos CPS médios se apresentaram bastante próximos entre cada vídeo. É importante ressaltar, que todas as legendas foram feitas com estratégia *verbatim*, ou seja, não houve redução nem reformulação textual, buscando seguir o estilo textual usado nos vídeos analisados previamente. A tabela abaixo mostra os valores médios de CPS obtidos para as legendas de cada vídeo:

**Tabela 5 – Valores de CPS médio para os vídeos preparados**

|                  | Vídeo 1 | Vídeo 2 | Vídeo 3 | Vídeo 4 | Vídeo 5 | Vídeo 6 | Média |
|------------------|---------|---------|---------|---------|---------|---------|-------|
| <b>CPS médio</b> | 14,92   | 14,75   | 14,27   | 14,50   | 14,28   | 14,46   | 14,53 |

Fonte: elaborado pelo autor.

Como pode ser observado dos valores acima, o valor médio de CPS ficou semelhante para todos os vídeos gravados, entre 14,27 e 14,95, demonstrando legendas na faixa de velocidades de lenta a média, como sugerem Nascimento (2018), Assis (2021), Vieira (2016) e Naves *et al.* (2016), considerando as tabelas propostas por Díaz Cintas e Remael (2021), discutidas na seção 2.2. Vale ressaltar que, em geral, os valores de CPS nos parâmetros de legendagem costumam incluir espaços em branco na contagem de caracteres. Contudo, assim como argumenta Rayner (1998), os leitores têm a tendência de não fixar o olhar em espaços em branco durante a leitura.

Em suma, os vídeos apresentaram uma média de 168 palavras, distribuídas durante uma gravação de 82 segundos de duração, cujas legendas foram criadas com média de CPS de 14,53. O tamanho médio das palavras foi de aproximadamente 7 letras por palavra apresentada, com frequência média de 4,34, o que representa uma frequência alta, segundo a escala zipf, tendo em vista que tanto o tamanho quanto a frequência das palavras são fatores que podem influenciar nos resultados, como apresentamos acima, na medida em que, em um experimento controlado, as fixações tendem a ser maiores em palavras menos frequentes (ver Rayner *et al.*, 2004), bem como em palavras com maior número de letras, ou seja, palavras com maior número de caracteres (Rayner; Sereno; Raney, 1996).

**Tabela 6** – Valores médios dos fatores de influência

| Vídeo | Duração (seg) | CPS AVG | Nº de Palavras | zipf_escala | nb_letras |
|-------|---------------|---------|----------------|-------------|-----------|
| Média | 82            | 14,53   | 168            | 4,345       | 7,123     |

Fonte: elaborado pelo autor

Com as legendas prontas, seguiu-se, então, para a escolha dos segmentos que seriam usados como textos em tela.

### **3.2.5 Delimitação dos textos em tela**

Para esta etapa, foram definidos, em cada vídeo, quatro pontos chaves a serem usados como os textos em tela. A fim de evitar interferência do número de

linhas no processamento linguístico, as legendas a serem usadas como texto em tela foram editadas em apenas uma linha. A seguir, é possível conferir o texto em tela de cada vídeo nas legendas, em negrito.

### 3.2.5.1 Legendas Manipuladas para o vídeo 1

**Tabela 7 – Texto em tela 1 para o vídeo 1**

| <b>Código de Tempo</b>            | <b>Duração (seg:frames)</b> | <b>Nº de caracteres</b> | <b>CPS</b> | <b>Legenda</b>                                 |
|-----------------------------------|-----------------------------|-------------------------|------------|------------------------------------------------|
| 00:00:04,931 --<br>> 00:00:07,821 | 02:27                       | 45                      | 15,57      | Aqui estão algumas dicas rápidas para melhorar |
| 00:00:07,845 --<br>> 00:00:09,646 | 01:24                       | 24                      | 13,33      | <b>a qualidade do seu sono:</b>                |
| 00:00:10,091 --<br>> 00:00:13,345 | 03:08                       | 44                      | 13,52      | Dica 1: Tenha uma Rotina de Sono Consistente:  |

Fonte: Elaborado pelo autor.

**Tabela 8 – Texto em tela 2 para o vídeo 1**

| <b>Códigos de Tempo</b>           | <b>Duração (seg:frames)</b> | <b>Nº de caracteres</b> | <b>CPS</b> | <b>Legenda</b>                   |
|-----------------------------------|-----------------------------|-------------------------|------------|----------------------------------|
| 00:00:29,478 --<br>> 00:00:30,507 | 01:01                       | 19                      | 18,46      | Seu quarto deve ser              |
| 00:00:30,531 --<br>> 00:00:31,921 | 01:12                       | 22                      | 15,83      | <b>um lugar calmo, escuro</b>    |
| 00:00:31,945 --<br>> 00:00:34,107 | 02:05                       | 32                      | 14,80      | e com uma temperatura agradável. |

Fonte: Elaborado pelo autor.

**Tabela 9 – Texto em tela 3 para o vídeo 1**

| <b>Códigos de Tempo</b>           | <b>Duração (seg:frames)</b> | <b>Nº de caracteres</b> | <b>CPS</b> | <b>Legenda</b>                               |
|-----------------------------------|-----------------------------|-------------------------|------------|----------------------------------------------|
| 00:00:39,849 --<br>> 00:00:41,787 | 01:28                       | 31                      | 16         | Opte por atividades relaxantes,              |
| 00:00:41,811 --<br>> 00:00:43,138 | 01:10                       | 18                      | 13,56      | <b>como ler um livro.</b>                    |
| 00:00:43,555 --<br>> 00:00:46,870 | 03:09                       | 43                      | 12,97      | Dica 3: Mantenha uma Alimentação balanceada: |

Fonte: Elaborado pelo autor.

**Tabela 10** – Texto em tela 4 para o vídeo 1

| <b>Códigos de Tempo</b>           | <b>Duração (seg:frames)</b> | <b>Nº de caracteres</b> | <b>CPS</b> | <b>Legenda</b>                                   |
|-----------------------------------|-----------------------------|-------------------------|------------|--------------------------------------------------|
| 00:01:02,145 --<br>> 00:01:04,999 | 02:26                       | 47                      | 14,47      | em uma grande melhoria na qualidade do seu sono. |
| 00:01:05,278 --<br>> 00:01:06,443 | 01:05                       | 11                      | 9,44       | <b>Comece hoje</b>                               |
| 00:01:06,467 --<br>> 00:01:08,000 | 01:16                       | 25                      | 16,31      | a implementar essas dicas                        |

Fonte: Elaborado pelo autor.

### 3.2.5.2 Legendas Manipuladas para o vídeo 2

**Tabela 11** – Texto em tela 1 para o vídeo 2

| <b>Códigos de Tempo</b>           | <b>Duração (seg:frames)</b> | <b>Nº de caracteres</b> | <b>CPS</b> | <b>Legenda</b>                                            |
|-----------------------------------|-----------------------------|-------------------------|------------|-----------------------------------------------------------|
| 00:00:09,943 --<br>> 00:00:13,818 | 03:21                       | 56                      | 14,45      | Por isso, separei três dicas essenciais para ter sucesso. |
| 00:00:14,703 --<br>> 00:00:16,532 | 01:20                       | 30                      | 16,40      | <b>Defina um horário de Trabalho:</b>                     |
| 00:00:16,876 --<br>> 00:00:19,155 | 02:07                       | 32                      | 14,04      | Estabeleça um horário consistente                         |

Fonte: Elaborado pelo autor.

**Tabela 12** – Texto em tela 2 para o vídeo 2

| <b>Códigos de Tempo</b>           | <b>Duração (seg:frames)</b> | <b>Nº de caracteres</b> | <b>CPS</b> | <b>Legenda</b>                                |
|-----------------------------------|-----------------------------|-------------------------|------------|-----------------------------------------------|
| 00:00:26,639 --<br>> 00:00:29,348 | 02:17                       | 44                      | 16,24      | e mantenha uma rotina que funcione para você. |
| 00:00:30,599 --<br>> 00:00:32,714 | 02:03                       | 27                      | 12,77      | <b>Vista-se de forma adequada:</b>            |
| 00:00:32,959 --<br>> 00:00:35,780 | 02:20                       | 40                      | 14,18      | Vestir-se como se estivesse no escritório     |

Fonte: Elaborado pelo autor.

**Tabela 13 – Texto em tela 3 para o vídeo 2**

| <b>Códigos de Tempo</b>           | <b>Duração (seg:frames)</b> | <b>Nº de caracteres</b> | <b>CPS</b> | <b>Legenda</b>                                          |
|-----------------------------------|-----------------------------|-------------------------|------------|---------------------------------------------------------|
| 00:00:53,193 --<br>> 00:00:54,749 | 01:13                       | 20                      | 12,85      | livre de distrações.                                    |
| 00:00:55,306 --<br>> 00:00:56,875 | 01:14                       | 26                      | 16,57      | <b>Tenha uma mesa organizada,</b>                       |
| 00:00:56,899 --<br>> 00:00:59,807 | 02:22                       | 51                      | 17,54      | uma cadeira confortável<br>e mantenha seu espaço limpo. |

Fonte: Elaborado pelo autor.

**Tabela 14 – Texto em tela 4 para o vídeo 2**

| <b>Códigos de Tempo</b>           | <b>Duração (seg:frames)</b> | <b>Nº de caracteres</b> | <b>CPS</b> | <b>Legenda</b>                                         |
|-----------------------------------|-----------------------------|-------------------------|------------|--------------------------------------------------------|
| 00:01:00,119 --<br>> 00:01:02,007 | 01:21                       | 27                      | 14,30      | Isso ajuda a manter o foco.                            |
| 00:01:02,784 --<br>> 00:01:04,406 | 01:15                       | 22                      | 13,56      | <b>Lembre-se dessas dicas</b>                          |
| 00:01:04,524 --<br>> 00:01:07,739 | 03:05                       | 50                      | 15,55      | para otimizar seu home office<br>e ser mais produtivo. |

Fonte: Elaborado pelo autor.

### 3.2.5.3 Legendas manipuladas para o vídeo 3

**Tabela 15 – Texto em tela 1 para o vídeo 3**

| <b>Códigos de Tempo</b>           | <b>Duração (seg:frames)</b> | <b>Nº de caracteres</b> | <b>CPS</b> | <b>Legenda</b>                 |
|-----------------------------------|-----------------------------|-------------------------|------------|--------------------------------|
| 00:00:10,785 --<br>> 00:00:11,785 | 01:00                       | 9                       | 9          | Primeiro:                      |
| 00:00:11,833 --<br>> 00:00:13,345 | 01:12                       | 22                      | 14,55      | <b>a desatenção constante</b>  |
| 00:00:13,370 --<br>> 00:00:15,133 | 01:18                       | 29                      | 16,45      | é uma característica marcante. |

Fonte: Elaborado pelo autor.

**Tabela 16 – Texto em tela 2 para o vídeo 3**

| <b>Códigos de Tempo</b>           | <b>Duração (seg:frames)</b> | <b>Nº de caracteres</b> | <b>CPS</b> | <b>Legenda</b>                       |
|-----------------------------------|-----------------------------|-------------------------|------------|--------------------------------------|
| 00:00:25,907 --<br>> 00:00:27,163 | 01:06                       | 17                      | 13,54      | Em segundo lugar,                    |
| 00:00:27,187 --<br>> 00:00:28,936 | 01:18                       | 29                      | 16,58      | <b>a hiperatividade e inquietude</b> |
| 00:00:28,960 --<br>> 00:00:29,963 | 01:00                       | 11                      | 10,97      | são comuns.                          |

Fonte: Elaborado pelo autor.

**Tabela 17 – Texto em tela 3 para o vídeo 3**

| <b>Códigos de Tempo</b>           | <b>Duração (seg:frames)</b> | <b>Nº de caracteres</b> | <b>CPS</b> | <b>Legenda</b>                    |
|-----------------------------------|-----------------------------|-------------------------|------------|-----------------------------------|
| 00:00:45,000 --<br>> 00:00:46,000 | 01:00                       | 11                      | 11         | Além disso,                       |
| 00:00:46,087 --<br>> 00:00:47,833 | 01:18                       | 26                      | 14,89      | <b>a impulsividade é notável.</b> |
| 00:00:47,993 --<br>> 00:00:49,533 | 01:13                       | 16                      | 10,39      | Pessoas com TDAH                  |

Fonte: Elaborado pelo autor.

**Tabela 18 – Texto em tela 4 para o vídeo 3**

| <b>Códigos de Tempo</b>           | <b>Duração (seg:frames)</b> | <b>Nº de caracteres</b> | <b>CPS</b> | <b>Legenda</b>                     |
|-----------------------------------|-----------------------------|-------------------------|------------|------------------------------------|
| 00:01:10,498 --<br>> 00:01:12,833 | 02:08                       | 26                      | 11,13      | O TDAH não é uma fraqueza.         |
| 00:01:12,903 --<br>> 00:01:14,433 | 01:13                       | 24                      | 15,69      | <b>É uma condição tratável.</b>    |
| 00:01:14,489 --<br>> 00:01:16,229 | 01:18                       | 33                      | 18,97      | Se você ou alguém que você conhece |

Fonte: Elaborado pelo autor.

### 3.2.5.4 Legendas manipuladas para o vídeo 4

**Tabela 19 – Texto em tela 1 para o vídeo 4**

| <b>Códigos de Tempo</b>           | <b>Duração (seg:frames)</b> | <b>Nº de caracteres</b> | <b>CPS</b> | <b>Legenda</b>                   |
|-----------------------------------|-----------------------------|-------------------------|------------|----------------------------------|
| 00:00:04,167 --<br>> 00:00:06,455 | 02:09                       | 32                      | 13,99      | requer algumas dicas essenciais. |
| 00:00:07,118 --<br>> 00:00:08,655 | 01:16                       | 19                      | 12,36      | <b>Treine até a falha.</b>       |
| 00:00:08,784 --<br>> 00:00:10,488 | 01:21                       | 31                      | 18,19      | Não importa se você é iniciante  |

Fonte: Elaborado pelo autor.

**Tabela 20 – Texto em tela 2 para o vídeo 4**

| <b>Códigos de Tempo</b>           | <b>Duração (seg:frames)</b> | <b>Nº de caracteres</b> | <b>CPS</b> | <b>Legenda</b>                      |
|-----------------------------------|-----------------------------|-------------------------|------------|-------------------------------------|
| 00:00:21,879 --<br>> 00:00:23,400 | 01:16                       | 23                      | 15,12      | e estimula o progresso.             |
| 00:00:24,039 --<br>> 00:00:26,021 | 01:29                       | 28                      | 14,13      | <b>Tenha a mentalidade correta:</b> |
| 00:00:26,211 --<br>> 00:00:28,155 | 01:28                       | 31                      | 15,95      | Desenvolver a mentalidade certa     |

Fonte: Elaborado pelo autor.

**Tabela 21 – Texto em tela 3 para o vídeo 4**

| <b>Códigos de Tempo</b>           | <b>Duração (seg:frames)</b> | <b>Nº de caracteres</b> | <b>CPS</b> | <b>Legenda</b>                     |
|-----------------------------------|-----------------------------|-------------------------|------------|------------------------------------|
| 00:00:59,379 --<br>> 00:01:00,921 | 01:16                       | 27                      | 17,51      | para um treinamento eficaz.        |
| 00:01:01,233 --<br>> 00:01:02,988 | 01:23                       | 27                      | 15,38      | <b>E lembre-se do aquecimento.</b> |
| 00:01:03,285 --<br>> 00:01:04,821 | 01:16                       | 23                      | 14,97      | Ele prepara os músculos            |

Fonte: Elaborado pelo autor.

**Tabela 22 – Texto em tela 4 para o vídeo 4**

| <b>Códigos de Tempo</b>           | <b>Duração (seg:frames)</b> | <b>Nº de caracteres</b> | <b>CPS</b> | <b>Legenda</b>                     |
|-----------------------------------|-----------------------------|-------------------------|------------|------------------------------------|
| 00:01:15,845 --<br>> 00:01:18,555 | 02:21                       | 29                      | 10,70      | nutrição, sono e aquecimento.      |
| 00:01:19,200 --<br>> 00:01:20,721 | 01:16                       | 22                      | 14,46      | <b>Essas dicas essenciais</b>      |
| 00:01:20,745 --<br>> 00:01:22,655 | 01:27                       | 33                      | 17,28      | garantem que você alcance o melhor |

Fonte: Elaborado pelo autor.

### 3.2.5.5 Legendas manipuladas para o vídeo 5

**Tabela 23 – Texto em tela 1 para o vídeo 5**

| <b>Códigos de Tempo</b>           | <b>Duração (seg:frames)</b> | <b>Nº de caracteres</b> | <b>CPS</b> | <b>Legenda</b>                      |
|-----------------------------------|-----------------------------|-------------------------|------------|-------------------------------------|
| 00:00:09,733 --<br>> 00:00:10,967 | 01:07                       | 15                      | 12,16      | Dica número um:                     |
| 00:00:11,132 --<br>> 00:00:13,155 | 02:01                       | 28                      | 13,84      | <b>Tenha uma dieta equilibrada,</b> |
| 00:00:13,300 --<br>> 00:00:15,088 | 01:24                       | 27                      | 15,10      | começando pela alimentação.         |

Fonte: Elaborado pelo autor.

**Tabela 24 – Texto em tela 2 para o vídeo 5**

| <b>Códigos de Tempo</b>           | <b>Duração (seg:frames)</b> | <b>Nº de caracteres</b> | <b>CPS</b> | <b>Legenda</b>              |
|-----------------------------------|-----------------------------|-------------------------|------------|-----------------------------|
| 00:00:29,052 --<br>> 00:00:30,255 | 01:06                       | 17                      | 14,13      | Dica número dois:           |
| 00:00:30,533 --<br>> 00:00:32,021 | 01:15                       | 20                      | 13,44      | <b>Pratique exercícios.</b> |
| 00:00:32,200 --<br>> 00:00:33,788 | 01:18                       | 25                      | 15,75      | Exercite-se regularmente,   |

Fonte: Elaborado pelo autor.

**Tabela 25 – Texto em tela 3 para o vídeo 5**

| <b>Códigos de Tempo</b>           | <b>Duração (seg:frames)</b> | <b>Nº de caracteres</b> | <b>CPS</b> | <b>Legenda</b>                                  |
|-----------------------------------|-----------------------------|-------------------------|------------|-------------------------------------------------|
| 00:00:48,300 --<br>> 00:00:49,588 | 01:09                       | 17                      | 13,20      | Dica número três:                               |
| 00:00:49,733 --<br>> 00:00:51,721 | 02:00                       | 26                      | 13,08      | <b>Priorize sua saúde mental.</b>               |
| 00:00:51,900 --<br>> 00:00:54,588 | 02:21                       | 46                      | 17,11      | Pratique técnicas de gerenciamento de estresse, |

Fonte: Elaborado pelo autor.

**Tabela 26 – Texto em tela 4 para o vídeo 5**

| <b>Códigos de Tempo</b>           | <b>Duração (seg:frames)</b> | <b>Nº de caracteres</b> | <b>CPS</b> | <b>Legenda</b>                     |
|-----------------------------------|-----------------------------|-------------------------|------------|------------------------------------|
| 00:01:10,733 --<br>> 00:01:11,767 | 01:01                       | 10                      | 9,67       | Lembre-se:                         |
| 00:01:11,779 --<br>> 00:01:13,655 | 01:26                       | 27                      | 14,39      | <b>Sua saúde é uma prioridade.</b> |
| 00:01:13,679 --<br>> 00:01:15,621 | 01:28                       | 32                      | 16,18      | Implemente essas mudanças simples  |

Fonte: Elaborado pelo autor.

### 3.2.5.6 Legendas manipuladas para o vídeo 6

**Tabela 27 – Texto em tela 1 para o vídeo 6**

| <b>Códigos de Tempo</b>           | <b>Duração (seg:frames)</b> | <b>Nº de caracteres</b> | <b>CPS</b> | <b>Legenda</b>                           |
|-----------------------------------|-----------------------------|-------------------------|------------|------------------------------------------|
| 00:00:12,933 --<br>> 00:00:14,121 | 01:06                       | 14                      | 11,78      | Primeira dica:                           |
| 00:00:14,145 --<br>> 00:00:16,388 | 02:07                       | 33                      | 14,71      | <b>Tenha uma alimentação balanceada.</b> |
| 00:00:16,412 --<br>> 00:00:18,788 | 02:11                       | 25                      | 10,52      | Consuma frutas, vegetais,                |

Fonte: Elaborado pelo autor.

**Tabela 28 – Texto em tela 2 para o vídeo 6**

| <b>Códigos de Tempo</b>           | <b>Duração (seg:frames)</b> | <b>Nº de caracteres</b> | <b>CPS</b> | <b>Legenda</b>                                        |
|-----------------------------------|-----------------------------|-------------------------|------------|-------------------------------------------------------|
| 00:00:28,933 --<br>> 00:00:30,161 | 01:07                       | 13                      | 10,59      | Segunda dica:                                         |
| 00:00:30,438 --<br>> 00:00:32,421 | 01:29                       | 32                      | 16,14      | <b>Alimente-se de forma consciente.</b>               |
| 00:00:32,592 --<br>> 00:00:36,188 | 03:18                       | 52                      | 14,46      | Evite distrações ao comer e concentre-se na refeição. |

Fonte: Elaborado pelo autor.

**Tabela 29 – Texto em tela 3 para o vídeo 6**

| <b>Códigos de Tempo</b>           | <b>Duração (seg:frames)</b> | <b>Nº de caracteres</b> | <b>CPS</b> | <b>Legenda</b>                                  |
|-----------------------------------|-----------------------------|-------------------------|------------|-------------------------------------------------|
| 00:00:43,000 --<br>> 00:00:44,155 | 01:05                       | 14                      | 12,12      | Terceira dica:                                  |
| 00:00:44,179 --<br>> 00:00:45,888 | 01:21                       | 26                      | 15,21      | <b>Hidrate-se com frequência.</b>               |
| 00:00:45,912 --<br>> 00:00:48,954 | 03:01                       | 46                      | 15,12      | Beba água regularmente para ajudar na digestão, |

Fonte: Elaborado pelo autor.

**Tabela 30 – Texto em tela 4 para o vídeo 6**

| <b>Códigos de Tempo</b>           | <b>Duração (seg:frames)</b> | <b>Nº de caracteres</b> | <b>CPS</b> | <b>Legenda</b>                            |
|-----------------------------------|-----------------------------|-------------------------|------------|-------------------------------------------|
| 00:01:15,345 --<br>> 00:01:17,488 | 02:04                       | 32                      | 14,93      | para orientações personalizadas.          |
| 00:01:17,733 --<br>> 00:01:19,655 | 01:28                       | 31                      | 16,13      | <b>Cuide de sua saúde e bem-estar,</b>    |
| 00:01:19,913 --<br>> 00:01:22,927 | 03:00                       | 40                      | 13,27      | fazendo escolhas alimentares conscientes. |

Fonte: Elaborado pelo autor.

### 3.2.6 Renderização dos Textos em Tela

Para a preparação dos estímulos, cada vídeo foi renderizado três vezes, sendo

estas:

1. Renderização dos vídeos contendo apenas os textos em tela;
2. Renderização dos vídeos com Legendas Totais;
3. Renderização dos vídeos com Legendas Parciais;

#### 3.2.6.1. Textos em tela

Cada texto em tela foi posicionado na região central do vídeo, assim como na maioria dos vídeos analisados na seção 2.2. A apresentação imagética desses elementos pode ser conferida nas imagens abaixo:

**Figura 23** – Representação imagética dos textos em tela utilizados.



Fonte: Elaborado pelo autor.

### 3.3 Condições Experimentais

Tendo em vista as discussões feitas até aqui, especialmente nas seções 3.4.1 e 3.4.2, foram estabelecidas duas condições experimentais de pesquisa, a saber:

- Condição 1: legendas totais
- Condição 2 legendas parciais

Para a renderização de cada vídeo, foram preparados dois arquivos de legenda em .srt: um incluindo todas as legendas do vídeo, a serem aplicadas nos vídeos categorizados na condição 1, e um segundo arquivo, cujas legendas redundantes nos segmentos com texto em tela foram excluídas e aplicadas nos vídeos categorizados na condição 2.

Observe, por exemplo, a sequência de legendas abaixo, retirada do Vídeo 1:

**Quadro 6 – Sequência de legendas para período crítico do vídeo 1.**

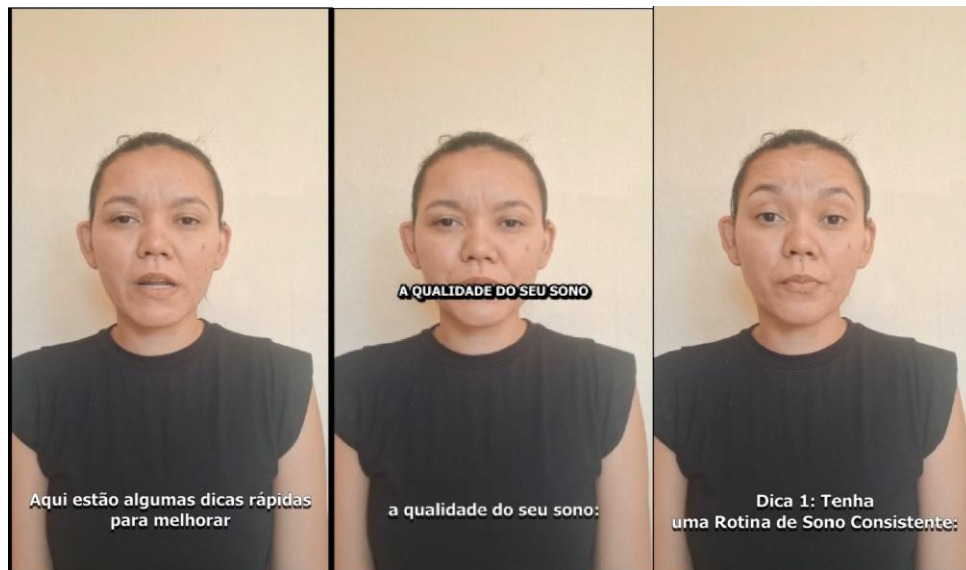
| Nº Legenda | Código de Tempo               | Texto da Legenda                               |
|------------|-------------------------------|------------------------------------------------|
| 3          | 00:00:04,931 --> 00:00:07,821 | Aqui estão algumas dicas rápidas para melhorar |
| 4          | 00:00:07,845 --> 00:00:09,646 | <b>a qualidade do seu sono:</b>                |
| 5          | 00:00:10,091 --> 00:00:13,345 | Dica 1: Tenha uma Rotina de Sono Consistente:  |

Fonte: elaborado pelo autor.

#### 3.3.1 Condição 1: Legendas Totais

No primeiro caso, legendas totais, não houve exclusão do texto de legenda, e ambas as informações são apresentadas visualmente, como ilustra a imagem abaixo.

**Figura 24** – Sequência de imagens para período crítico no vídeo 1 na condição total.

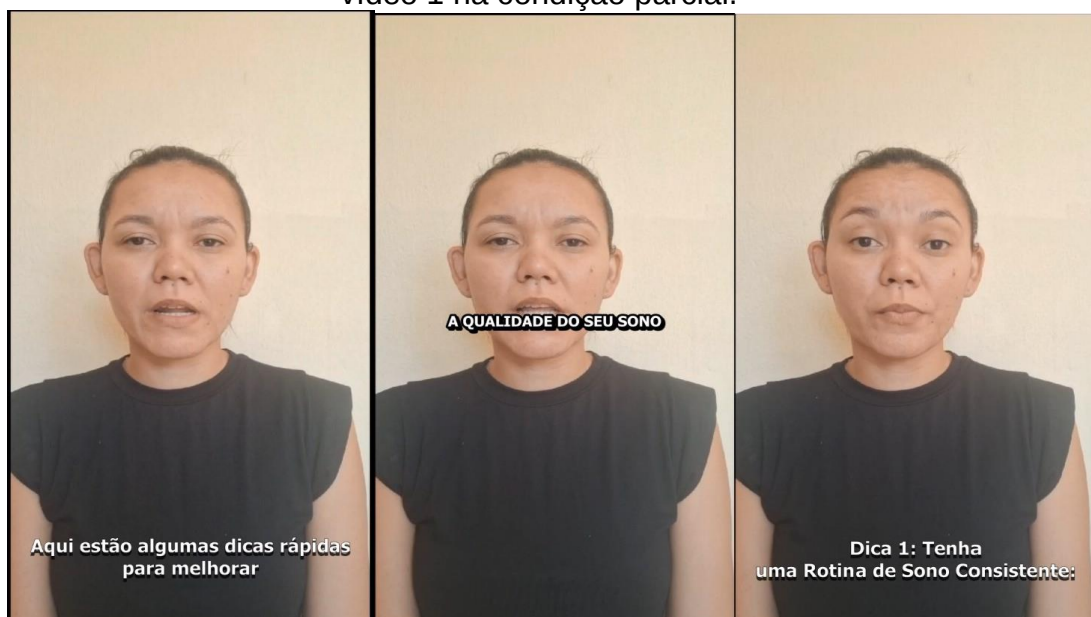


Fonte: elaborado pelo autor.

### 3.3.2 Condição 1: Legendas Parciais

Já no segundo caso, aparecem as legendas parciais, onde houve a exclusão da legenda referente ao segmento do texto que foi usado como texto em tela.

**Figura 25** – Captura de Tela: Sequência de imagens para período crítico no vídeo 1 na condição parcial.



Fonte: elaborado pelo autor

O mesmo processo foi feito para os demais vídeos, e, para cada área de interesse, delimitou-se uma legenda antes e uma legenda após a aparição do texto em tela, tendo em vista que essas seriam as responsáveis pela migração da atenção e transição entre um texto e outro.

### **3.4 Experiment Builder**

Para a montagem do experimento, foi usado o *Experiment Builder*, do *Eye Link 1000 Plus* – 2000 Hz de frequência – (SR RESEARCH, 2022), que permite fazer o *deploy* de uma versão executável do material. Por questões de compatibilidade com a resolução do *Eye Link* e execução no *Experiment Builder*, todos os vídeos foram convertidos para 16x9 (*ratio* horizontal), mantendo, assim mesmo, a posição vertical dos vídeos principais do experimento, com taxa de *frame* de 30 FPS (*frames* por segundo) e a resolução de 1280x720. O áudio do material foi retirado a fim de não obter interferência desse segmento, como pontua Godfroid (2019, p. 165), já que “em um experimento típico, espera-se que os participantes direcionem a atenção para um objeto na tela quando o rótulo do objeto é mencionado no áudio”, ou seja, quando esse objeto que está sendo visualizado é mencionado/nomeado no áudio.

Para garantir que a ordem de exibição dos vídeos não influenciasse os resultados (Orero *et al*, 2018), cada participante assistiu a uma lista de vídeos contendo 6 vídeos principais – 3 com legendas totais e 3 com legendas parciais –, distribuídos em 4 listas aleatórias (A, B, C e D), acrescidos de 12 distratores – vídeos retirados do YouTube e Instagram com legendas totais e parciais – com durações entre 60 e 90 segundos. Os distratores, além de conterem vídeos variados com ambas as condições de pesquisa, funcionam como um sinal de alerta facilitador, como pontua Holmqvist *et al* (2011), sem entregar ao participante o foco de observação do material para enviesar o comportamento de leitura e comprometer os dados da coleta.

Sendo assim, cada lista foi composta de 18 vídeos em exibição aleatória. A tabela abaixo exemplifica a ordem de exibição dos vídeos para a Lista A. Todas as listas podem ser conferidas aqui.

**Tabela 31 – Lista aleatória A**

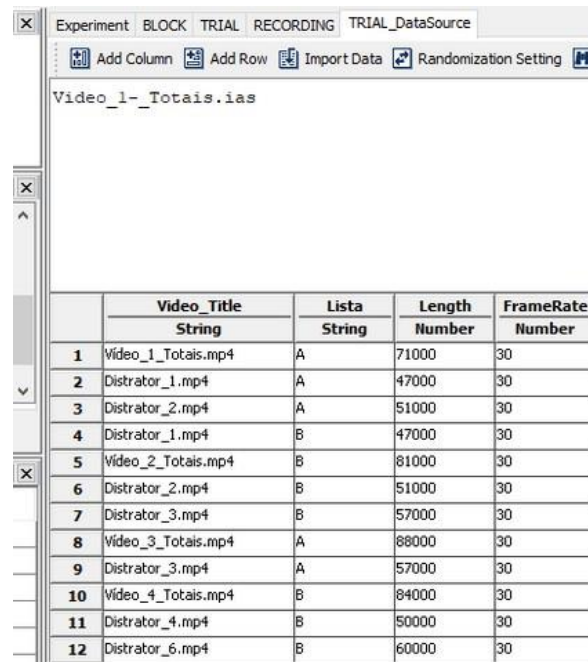
| <b>Título</b>   | <b>Resolução</b> | <b>Taxa de quadros</b> |
|-----------------|------------------|------------------------|
| Video 1_Total   | 1280 x 720       | 30                     |
| Distrator 1     | 1280 x 720       | 30                     |
| Distrator 2     | 1280 x 720       | 30                     |
| Video 3_Total   | 1280 x 720       | 30                     |
| Distrator 3     | 1280 x 720       | 30                     |
| Video 2_Parcial | 1280 x 720       | 30                     |
| Distrator 4     | 1280 x 720       | 30                     |
| Distrator 10    | 1280 x 720       | 30                     |
| Distrator 5     | 1280 x 720       | 30                     |
| Video 5_Total   | 1280 x 720       | 30                     |
| Distrator 6     | 1280 x 720       | 30                     |
| Video 4_Parcial | 1280 x 720       | 30                     |
| Distrator 7     | 1280 x 720       | 30                     |
| Distrator 8     | 1280 x 720       | 30                     |
| Distrator 9     | 1280 x 720       | 30                     |
| Video 6_Parcial | 1280 x 720       | 30                     |
| Distrator 11    | 1280 x 720       | 30                     |
| Distrator 12    | 1280 x 720       | 30                     |
| Distrator 1     | 1280 x 720       | 30                     |
| Video 2_Total   | 1280 x 720       | 30                     |
| Distrator 2     | 1280 x 720       | 30                     |
| Distrator 3     | 1280 x 720       | 30                     |
| Video 4_Total   | 1280 x 720       | 30                     |
| Distrator 4     | 1280 x 720       | 30                     |
| Distrator 6     | 1280 x 720       | 30                     |
| Video 1_Parcial | 1280 x 720       | 30                     |
| Distrator 7     | 1280 x 720       | 30                     |
| Video 6_Total   | 1280 x 720       | 30                     |

|                  |            |    |
|------------------|------------|----|
| Distrator 5      | 1280 x 720 | 30 |
| Distrator 9      | 1280 x 720 | 30 |
| Distrator 10     | 1280 x 720 | 30 |
| Video 3_Parcial  | 1280 x 720 | 30 |
| Video 5_Parcial  | 1280 x 720 | 30 |
| Distrator 11     | 1280 x 720 | 30 |
| Distrator 8      | 1280 x 720 | 30 |
| Distrator 12     | 1280 x 720 | 30 |
| Distrator_4      | 1280 x 720 | 30 |
| Distrator_3      | 1280 x 720 | 30 |
| Distrator_1      | 1280 x 720 | 30 |
| Video_2_Totais   | 1280 x 720 | 30 |
| Distrator_10     | 1280 x 720 | 30 |
| Video_5_Parciais | 1280 x 720 | 30 |
| Distrator_2      | 1280 x 720 | 30 |
| Distrator_5      | 1280 x 720 | 30 |
| Distrator_6      | 1280 x 720 | 30 |
| Video_3_Parciais | 1280 x 720 | 30 |
| Distrator_12     | 1280 x 720 | 30 |
| Vídeo_4_Totais   | 1280 x 720 | 30 |
| Distrator_8      | 1280 x 720 | 30 |
| Video_1_Parciais | 1280 x 720 | 30 |
| Distrator_7      | 1280 x 720 | 30 |
| Video_6_Totais   | 1280 x 720 | 30 |
| Distrator_11     | 1280 x 720 | 30 |
| Distrator_9      | 1280 x 720 | 30 |
| Video_5_Totais   | 1280 x 720 | 30 |
| Distrator_1      | 1280 x 720 | 30 |
| Distrator_3      | 1280 x 720 | 30 |
| Video_2_Parciais | 1280 x 720 | 30 |

|                  |            |    |
|------------------|------------|----|
| Distrator_10     | 1280 x 720 | 30 |
| Distrator_4      | 1280 x 720 | 30 |
| Distrator_2      | 1280 x 720 | 30 |
| Distrator_5      | 1280 x 720 | 30 |
| Distrator_12     | 1280 x 720 | 30 |
| Video_4_Totais   | 1280 x 720 | 30 |
| Distrator_8      | 1280 x 720 | 30 |
| Video_3_Totais   | 1280 x 720 | 30 |
| Distrator_6      | 1280 x 720 | 30 |
| Distrator_9      | 1280 x 720 | 30 |
| Video_1_Parciais | 1280 x 720 | 30 |
| Distrator_7      | 1280 x 720 | 30 |
| Video_6_Parciais | 1280 x 720 | 30 |
| Distrator_11     | 1280 x 720 | 30 |

Fonte: elaborado pelo autor.

Uma vez que todo o material estava pronto – os valores de duração dos vídeos, listas pertencentes, taxa de quadros e resolução –, essas informações foram transferidas para o “TRIAL\_datasource”, no *Experiment Builder*, como ilustra a Fig. 26, abaixo.

**Figura 26 – Interface do TRIAL\_Datasource no Experiment Builder**


|    | Video Title<br>String | Lista<br>String | Length<br>Number | FrameRate<br>Number |
|----|-----------------------|-----------------|------------------|---------------------|
| 1  | Video_1_Totais.mp4    | A               | 71000            | 30                  |
| 2  | Distrator_1.mp4       | A               | 47000            | 30                  |
| 3  | Distrator_2.mp4       | A               | 51000            | 30                  |
| 4  | Distrator_1.mp4       | B               | 47000            | 30                  |
| 5  | Video_2_Totais.mp4    | B               | 81000            | 30                  |
| 6  | Distrator_2.mp4       | B               | 51000            | 30                  |
| 7  | Distrator_3.mp4       | B               | 57000            | 30                  |
| 8  | Video_3_Totais.mp4    | A               | 88000            | 30                  |
| 9  | Distrator_3.mp4       | A               | 57000            | 30                  |
| 10 | Video_4_Totais.mp4    | B               | 84000            | 30                  |
| 11 | Distrator_4.mp4       | B               | 50000            | 30                  |
| 12 | Distrator_6.mp4       | B               | 60000            | 30                  |

Fonte: elaborado pelo autor.

É possível delimitar as áreas de interesse no próprio *Experiment Builder*, optar por defini-las e aplicá-las no *Data Viewer* durante o processo de análise de dados e/ou importá-las ao *Experiment Builder*, para que sejam aplicadas de forma automática a cada participante/trial durante a coleta.

Em uma primeira versão do experimento, havia sido coletada uma amostra teste, cujas áreas de interesse foram incorporadas à montagem do experimento. No entanto, devido a uma atualização dos métodos de análise, optou-se por redefini-las e aplicá-las diretamente no *Data Viewer*, de acordo com os objetivos de pesquisa, assim como descrito nas seções 6.1 e 6.2.

### 3.5 Áreas de Interesse

Tendo em vista a disponibilidade limitada de trabalhos cuja metodologia se destinasse à observação de vídeos com legendas e pontos estratégicos do vídeo, optou-se por adaptar as metodologias existentes na delimitação de áreas de interesse para materiais audiovisuais.

A maioria dos estudos sobre rastreamento ocular destina-se à análise da leitura em textos estáticos. No caso dos vídeos legendados, temos uma situação

envolvendo textos dinâmicos, o que dificulta o processo de delimitação de áreas de interesse como nas metodologias usuais.

Assim como descrito na seção teórica, tanto nos estudos de Liao et al (2021) e Assis (2021), os autores optaram por delimitar áreas de interesse com abrangência de toda a área de legenda. Liao, Kruger e Doherty (2020) também usam a totalidade da área do vídeo como uma área de interesse, investigando a divisão de atenção entre área de legenda e área de vídeo. Josephson e Holmes (2006) também optam por definir áreas de interesse abrangentes para a área do vídeo a que se destina a análise. Contudo, não se replica essa AI no presente trabalho, tendo em vista, inicialmente, o objetivo isolado da pesquisa.

Em resumo, os estudos que mais se aproximam do presente no que tange à divisão das áreas de interesse são Liao et al (2021) e Liao, Kruger e Doherty (2020). Aqui, propõe-se, então, uma adaptação das duas pesquisas às condições experimentais já descritas.

- Liao et al (2021): enquanto optam por utilizar legendas com ou sem a presença de vídeo; aqui, opta-se por utilizar vídeos com legendas totais ou parciais.
- Liao, Kruger e Doherty (2020): avaliam a divisão de atenção nas legendas bilíngues (chinês e inglês), com duas áreas de interesse principais (uma para cada legenda); aqui, também se opta por utilizar duas áreas de interesse, a saber: AI 1 – correspondendo à área das legendas, na parte inferior dos vídeos – e AI 2 – correspondendo à área dos textos em tela, na parte central dos vídeos.

### **3.6 Aspectos Éticos da Pesquisa**

A presente pesquisa foi executada mediante envio do processo de número 80810724.0.0000.5054 ao Comitê de Ética da Universidade Federal do Ceará/PROPESQ, através da Plataforma Brasil.

Todos os participantes foram informados sobre os objetivos e a metodologia da pesquisa e convidados a lerem e assinarem o Termo de Consentimento Livre e Esclarecido (TCLE), disponível no Apêndice 10, concordando em participar da pesquisa.

### **3.7 Etapas de Coleta**

Esta seção destina-se a apresentar o passo a passo proposto para a coleta de dados do piloto do experimento descrito acima, a qual se constitui de três etapas:

- Etapa 1: Preenchimento de Questionário Pré-Coleta (Apêndice A), com perguntas voltadas ao perfil socioeconômico dos sujeitos, bem como de seus hábitos de uso nas redes sociais, especialmente, em relação ao consumo de vídeos legendados, seguido pela assinatura do TCLE, quando os participantes eram informados sobre todos os procedimentos de coleta, bem como os riscos e benefícios em participar da pesquisa;
- Etapa 2: Execução do material no *Eye Link*, com gravação dos estímulos;
- Etapa 3: Resposta ao questionário Pós-Coleta (Apêndice B), com perguntas sobre o conteúdo dos vídeos, a fim de se analisar o nível de entendimento dos participantes sobre o material apresentado.

#### **3.7.1 Procedimentos Pré-coleta**

Para auxiliar na coleta de dados e organização dos materiais, bem como dos questionários, optou-se por usar a *Airtable* (2023), uma plataforma multitarefas, que permite ao usuário personalizar o uso e interface de acordo com seus objetivos. A plataforma conta com armazenamento de dados em nuvem, integrações e automações, que possibilitam a criação de questionários e fórmulas a fim de facilitar o processo de coleta, armazenamento e análise de dados.

Para os objetivos desta pesquisa, a *Airtable* foi usada para:

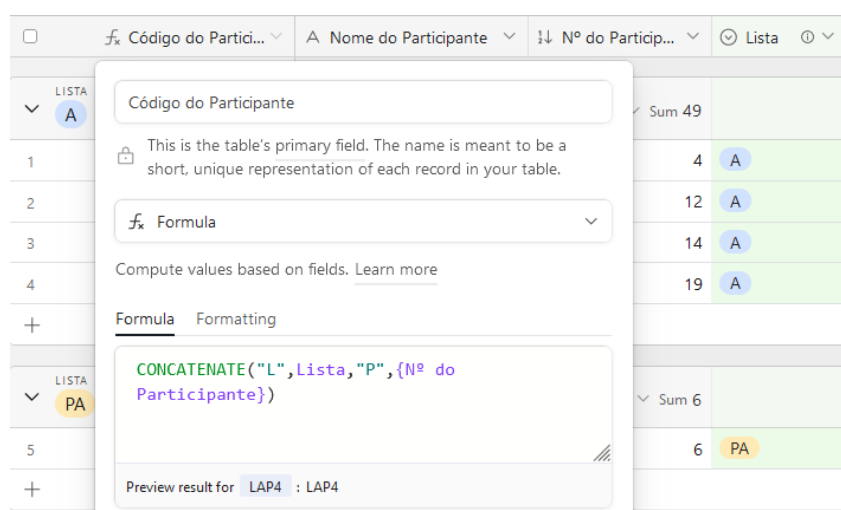
- Listagem de Palavras dos Roteiros (ver Apêndices C e H)
- Criação dos Questionários Pré e Pós-Coleta;
- Geração dos Códigos dos Participantes.

#### **3.7.2 Procedimento de Coleta**

Na etapa 1, a lista de exibição é atribuída ao participante no momento de responder ao questionário Pré-Coleta, e, a fim de manter o anonimato na listagem

de participantes, cada participante recebe um código de usuário gerado através de uma fórmula padrão na *Airtable* (2023). O código é gerado pela fórmula ("L",Lista,"P",{Nº do Participante}), e o número do participante não pode ser alterado, uma vez que o questionário é respondido, o que garante maior precisão na listagem de forma anônima. Sendo assim, ao participante 4, por exemplo, foi atribuído o código LAP4 – o que significa que o Participante 4 viu a lista A –, que também foi atribuído no *Experiment Builder*, durante a etapa 2.

**Figura 27 – Criação de fórmula para códigos dos participantes.**



Fonte: elaborado pelo autor.

### 3.7.3 Participantes para a coleta

Dos 18 participantes que aceitaram participar do piloto, apenas oito puderam comparecer ao laboratório para realizar a coleta de dados. O principal período de coleta ocorreu durante a greve dos servidores, o que impactou no número baixo de participantes. Desses oito, os dados de dois participantes precisaram ser descartados (participantes LAP19 e LAP14). O primeiro, devido à baixa precisão durante a calibração; o segundo, por ter tirado a cabeça do suporte durante o experimento e ter tido distrações em outros *trials* devido ao toque do celular, que não foi deixado no modo silencioso. Portanto, no piloto, foram obtidos dados de seis participantes válidos para a análise.

A seguir, apresenta-se o perfil geral dos sujeitos, com base no questionário pré-coleta a que responderam:

**Quadro 7 – Perfil Socioeconômico dos participantes em idade, sexo e nível de escolaridade**

| Participante | Idade | Sexo | Escolaridade                        |
|--------------|-------|------|-------------------------------------|
| LPAP6        | 29    | F    | Pós-Graduação (Mestrado) Completo   |
| LDP11        | 25    | F    | Ensino Superior Incompleto          |
| LDP16        | 22    | F    | Ensino Superior Completo            |
| LDP17        | 29    | M    | Ensino Superior Incompleto          |
| LDP20        | 24    | M    | Ensino Superior Incompleto          |
| LDP23        | 24    | M    | Pós-Graduação (Mestrado) Incompleto |

Fonte: elaborado pelo autor.

Os participantes tinham entre 22 e 29 anos de idade (25,5 em média). Três participantes eram do sexo masculino, e três do sexo feminino. Em relação ao nível de escolaridade, dois participantes foram alunos de pós-graduação, e quatro alunos de graduação.

No que tange ao perfil dos sujeitos em relação ao consumo de conteúdos audiovisuais legendados, três sujeitos (LPAP6, LDP11 e LDP16) afirmaram não ter costume de ver filmes ou programas de TV legendados. Ao responderem se tinham alguma dificuldade ao ver conteúdo legendado, LPAP6, LDP16, LDP17 afirmaram que sim, especificando qual problema:

LPAP6: *“Não assisto filme, pois não gosto. Quando assisto documentário, não gosto de legendas que aparecem aos poucos (por letras ou sílaba)”.*

LDP16: *“Concentração”*

LDP17: *“É difícil acompanhar a legenda e as imagens, quando foco em um, perco informações do outro.”*

Em relação ao perfil de consumo de conteúdo nas redes sociais, todos os participantes afirmaram acessar as redes sociais. Com exceção de LPAP6, que, embora acesse, o faz “raramente”; todos os demais disseram utilizar diariamente.

tempo de acesso às redes sociais variou de “menos de 1h por dia”, para LPAP6, a mais de 5 horas por dia, para LDP11 e LDP20. As principais redes sociais listadas foram “Instagram” e “YouTube”, apesar de serem listadas também TikTok, Twitter e Facebook.

**Quadro 8** – Perfil dos participantes em relação aos hábitos de acesso a redes sociais

| Participante | Usa Redes Sociais | Frequência  | Horas/dia     | Redes Sociais   |
|--------------|-------------------|-------------|---------------|-----------------|
| LPAP6        | Sim               | Raramente   | < 1h          | IG e YT         |
| LDP11        | Sim               | Diariamente | > 5h          | IG e YT         |
| LDP16        | Sim               | Diariamente | Entre 1h e 3h | IG, YT, TK      |
| LDP17        | Sim               | Diariamente | Entre 3h e 5h | IG, TK, Twitter |
| LDP20        | Sim               | Diariamente | > 5h          | IG, FB, TK      |
| LDP23        | Sim               | Diariamente | Entre 3h e 5h | IG e YT         |

Fonte: elaborado pelo autor

### 3.7.3.1 Listas assistidas

Assim como mencionado na seção 5.4.2, aos participantes, foram atribuídas listas aleatórias durante a resposta aos questionários. Como nem todos os que aceitaram participar concluíram a coleta, cinco dos participantes assistiram à lista D e um participante assistiu à lista A. Os vídeos presentes em cada lista são especificados a seguir:

**Quadro 9** – Sequência de vídeos e suas condições experimentais por lista assistida

| Participante | Lista | Trial | Vídeo/Condição   |
|--------------|-------|-------|------------------|
| LPAP6        | A     | 1     | Video 1_Totais   |
|              |       | 4     | Video 3_Totais   |
|              |       | 6     | Video 2_Parciais |
|              |       | 10    | Video 5_Totais   |

|       |   |    |                  |
|-------|---|----|------------------|
|       |   | 12 | Video 4_Parciais |
|       |   | 16 | Video 6_Parciais |
| LDP11 | D | 1  | Vídeo_5_Totais   |
| LDP16 |   | 4  | Vídeo_2_Parciais |
| LDP17 |   | 10 | Vídeo_4_Totais   |
| LDP20 |   | 12 | Vídeo_3_Totais   |
| LDP23 |   | 15 | Vídeo_1_Parciais |
|       |   | 17 | Vídeo_6_Parciais |

Fonte: elaborado pelo autor.

Todos os participantes assistiram aos mesmos vídeos, porém, em condições diferentes, sendo estes: três vídeos em condição de legenda total e três vídeos em condição de legenda parcial.

Na seção seguinte, são apresentados os resultados preliminares obtidos quanto ao percurso de análise dos dados dentro dos softwares do *Eye Link*.

### 3.8.1 Tratamento dos dados da movimentação ocular

Os dados da coleta foram processados no *Data Viewer*, versão 4.1.1, que é uma das ferramentas do *Eye Link*. Através do *Data Viewer*, é possível visualizar as gravações do movimento ocular de cada participante, em vídeo, imagem e planilhas, com relatórios de áreas de interesse, sacadas e fixações, por exemplo.

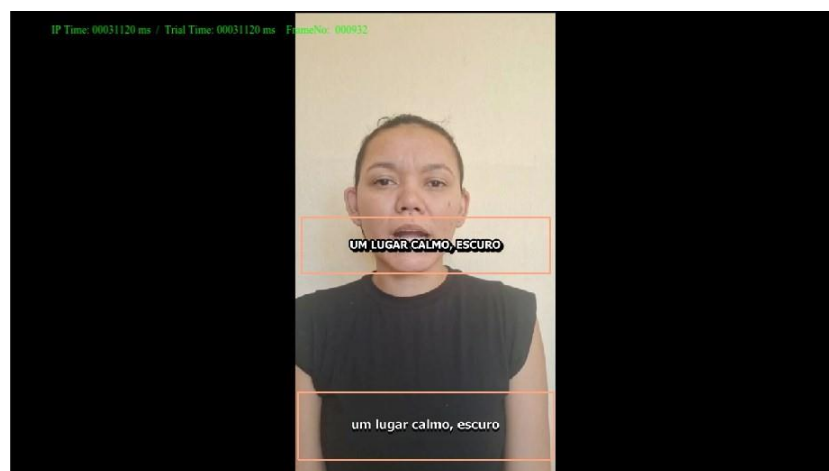
- **Áreas de Interesse**

As áreas de interesse foram definidas no próprio *Data Viewer* para cada vídeo e aplicadas a cada participante. Cada vídeo teve duas áreas de interesse, tabeladas da seguinte forma:

- AI 1 - “*bottom*”: considerando a região das legendas na parte inferior dos vídeos;
- AI 2 - “*middle*”, considerando a região dos textos em tela na parte central dos vídeos.

Não se objetiva, inicialmente, verificar o padrão de movimento nas demais áreas do vídeo, já que o objeto de estudo principal é o uso de legendas na presença dos textos em tela.

**Figura 28** – Áreas de interesse delimitadas para os vídeos



Fonte: elaborado pelo autor.

Os tempos de entrada e saída de cada AI foram definidos usando o tempo IP do experimento para a duração do vídeo em questão. Na Fig. 28, acima, por exemplo, é possível ter acesso ao “IP Time” do *frame*.

Para definir o “IP Time” de cada AI, inicialmente, foi definido um período de interesse principal, nomeado como “fullvideos”, considerando o tempo total do experimento. Essa etapa foi feita no “*Interest Period Manager*”. Para definir o tempo inicial do IP, foi aplicada a mensagem “*Frame to be displayed 1*”. Para definir o tempo final do IP, foi usada a mensagem “*display\_blank*”. Ambos os eventos foram programados na montagem do experimento, como recomendado pelo suporte, bem como definido no Manual do Usuário (Eyelink, 2024, p. 119). A Fig. 29, abaixo, ilustra as mensagens usadas na criação do período de interesse principal.

**Figura 29** – Período de Interesse Principal, “fullvideos”, no “Interest Period Editor”

The screenshot shows the 'Interest Period Editor' window. At the top, the 'Interest Period Name' is set to 'fullvideos'. Below this, the 'IP Start Event Settings' section is visible, with 'Interest Period Start Event Type' set to 'EDF Message' and 'Message Text' set to 'Frame to be display'. The 'IP End Event Settings' section is also visible, with 'Interest Period End Event Type' set to 'EDF Message' and 'Message Text' set to 'display\_blank'. The 'General Options' section at the bottom has 'Apply Strict Event Matching' checked. The window includes 'Reset', 'Cancel', and 'OK' buttons at the bottom right.

Fonte: elaborado pelo autor.

No entanto, como não se objetivou analisar o movimento ocular em todo o vídeo, foram definidos períodos de interesse para cada texto em tela, ou seja, um total de quatro períodos de interesse para cada vídeo, os quais foram nomeados como “CritPeriod[número]”. E, para cada período crítico, foram adicionadas mensagens com os tempos de entrada e saída de cada texto em tela e sua respectiva legenda (no caso da condição total). Esse processo é descrito na seção seguinte.

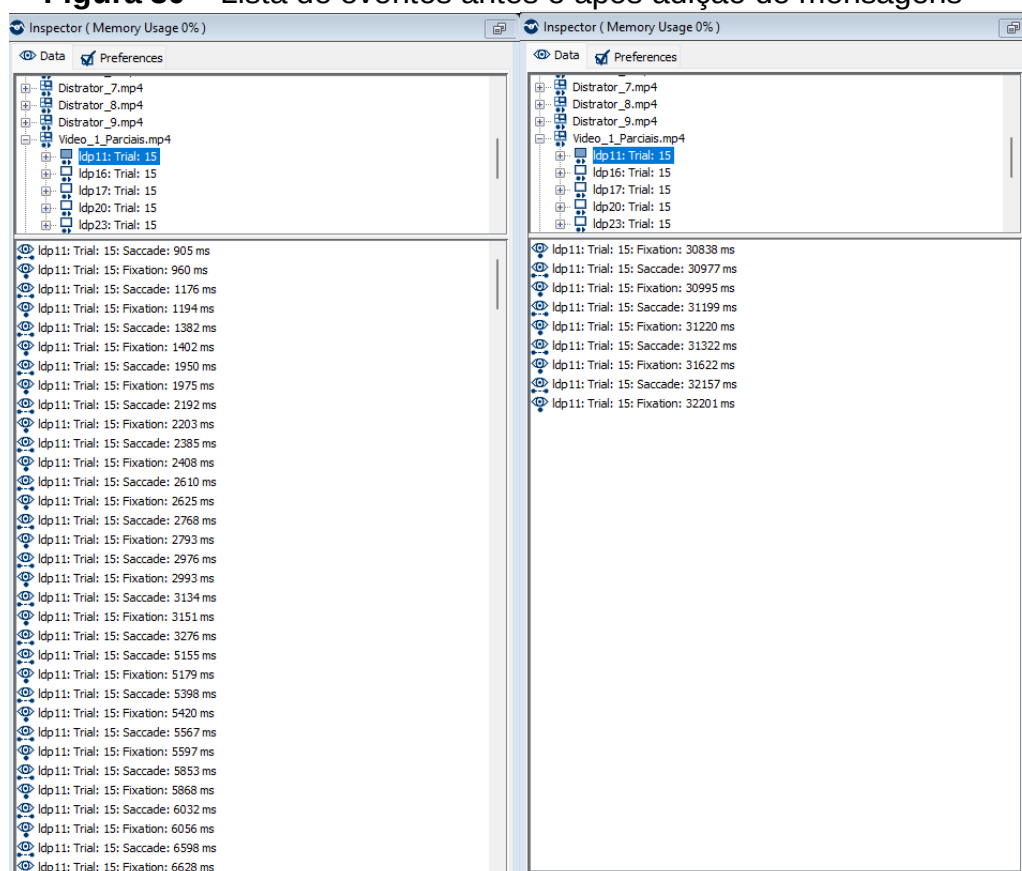
### 3.8.1.2. *Períodos de Interesse*

Cada gravação envia uma mensagem de acordo com cada evento, seja este fixações, piscadelas, sacadas etc. No entanto, ao se considerar todas as mensagens, durante a análise, também se obtém todos esses valores de acordo com cada relatório.

Sendo assim, precisaram ser definidas mensagens específicas que identificassem o tempo de entrada e saída de cada período de interesse, apenas nos segmentos do vídeo onde havia a presença dos textos em tela, mais

especificamente, como mencionado anteriormente, em quatro momentos para cada vídeo. Esse passo garante que, ao exportar os relatórios de análise e dados de imagem, apenas sejam incluídos os segmentos que se pretende analisar. Na Fig. 30, abaixo, é possível ver, à esquerda, toda a lista de mensagens e eventos enviados durante a coleta de dados de um dos participantes, sem a aplicação dos períodos de interesse; à direita, a mesma lista de mensagens, porém, com a aplicação dos períodos de interesse.

**Figura 30 – Lista de eventos antes e após adição de mensagens**

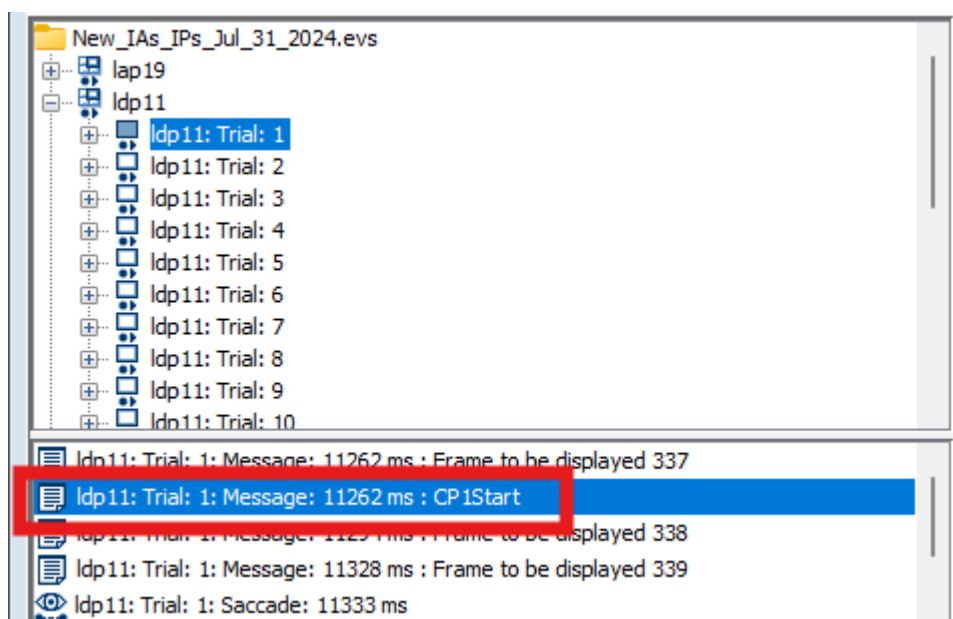


Fonte: elaborado pelo autor.

Para definir as mensagens, inicialmente, foi usado o Período de Interesse principal, descrito na seção anterior (ver seção 3.8.1), contendo todas as mensagens e eventos de cada *trial*. Cada mensagem de entrada foi definida como CP + [número no período de interesse] + [Start]. De forma semelhante, as mensagens de saída foram definidas com a mesma titulação, mas adicionado [End], para indicar o tempo de saída das legendas. O CP foi definido como “período crítico”, para diferenciar a delimitação dos tempos de entrada e saída de

cada mensagem e não confundir com o “Período de Interesse” definido no “*Interest Period Manager*”. Sendo assim, a mensagem para o Período Crítico de entrada 1 para o *trial* 1, por exemplo, foi definida como “CP1Start” e é mostrada na lista de mensagens dos eventos, como mostra a figura abaixo:

**Figura 31** – Exemplo de código de mensagem na lista de eventos



Fonte: elaborado pelo autor.

O mesmo procedimento foi repetido para cada um dos períodos críticos de cada vídeo e aplicados a cada participante, de acordo com cada *trial*. Sendo assim, a lista de mensagens do participante LDP11, por exemplo, resultou nos dados exemplificados no quadro 10, abaixo:

**Quadro 10** – Exemplo de lista de mensagem para a lista de visualização D

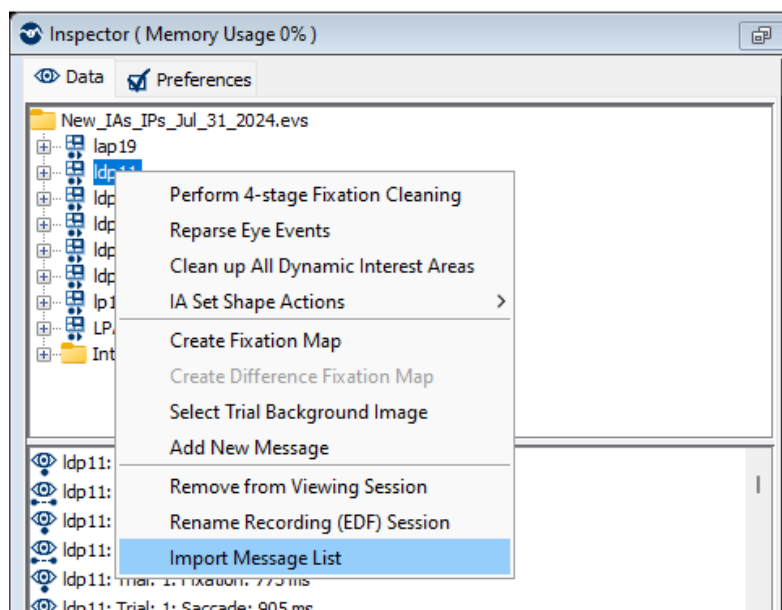
| Mensagem ID | Trial | Mensagem - IP time | Mensagem - Text ID |
|-------------|-------|--------------------|--------------------|
| MSG         | -1    | -11199             | CP1Start           |
| MSG         | -1    | -13332             | CP1End             |
| MSG         | -1    | -30697             | CP2Start           |
| MSG         | -1    | -32264             | CP2End             |
| MSG         | -1    | -49729             | CP3Start           |
| MSG         | -1    | -51929             | CP3End             |
| MSG         | -1    | -71860             | CP4Start           |
| MSG         | -1    | -73726             | CP4End             |
| MSG         | -4    | -14888             | CP1Start           |
| MSG         | -4    | -16599             | CP1End             |
| MSG         | -4    | -30791             | CP2Start           |
| MSG         | -4    | -32948             | CP2End             |
| MSG         | -4    | -55462             | CP3Start           |
| MSG         | -4    | -56928             | CP3End             |
| MSG         | -4    | -62962             | CP4Start           |
| MSG         | -4    | -64562             | CP4End             |
| MSG         | -10   | -7299              | CP1Start           |
| MSG         | -10   | -8832              | CP1End             |
| MSG         | -10   | -24098             | CP2Start           |
| MSG         | -10   | -26164             | CP2End             |
| MSG         | -10   | -61261             | CP3Start           |
| MSG         | -10   | -63228             | CP3End             |
| MSG         | -10   | -79227             | CP4Start           |
| MSG         | -10   | -80827             | CP4End             |
| MSG         | -12   | -12000             | CP1Start           |
| MSG         | -12   | -13366             | CP1End             |
| MSG         | -12   | -27198             | CP2Start           |
| MSG         | -12   | -29131             | CP2End             |
| MSG         | -12   | -46230             | CP3Start           |

|     |     |        |          |
|-----|-----|--------|----------|
| MSG | -12 | -48030 | CP3End   |
| MSG | -12 | -72944 | CP4Start |
| MSG | -12 | -74494 | CP4End   |
| MSG | -15 | -8213  | CP1Start |
| MSG | -15 | -10319 | CP1End   |
| MSG | -15 | -30771 | CP2Start |
| MSG | -15 | -32134 | CP2End   |
| MSG | -15 | -42028 | CP3Start |
| MSG | -15 | -43555 | CP3End   |
| MSG | -15 | -65599 | CP4Start |
| MSG | -15 | -66661 | CP4End   |
| MSG | -17 | -14888 | CP1Start |
| MSG | -17 | -16599 | CP1End   |
| MSG | -17 | -30791 | CP2Start |
| MSG | -17 | -32948 | CP2End   |
| MSG | -17 | -44196 | CP3Start |
| MSG | -17 | -45862 | CP3End   |
| MSG | -17 | -77760 | CP4Start |
| MSG | -17 | -79660 | CP4End   |

Fonte: Elaborado pelo autor.

Cada arquivo de mensagem foi salvo em formato .txt para cada participante e importado para o *Data Viewer*, como ilustra a figura abaixo:

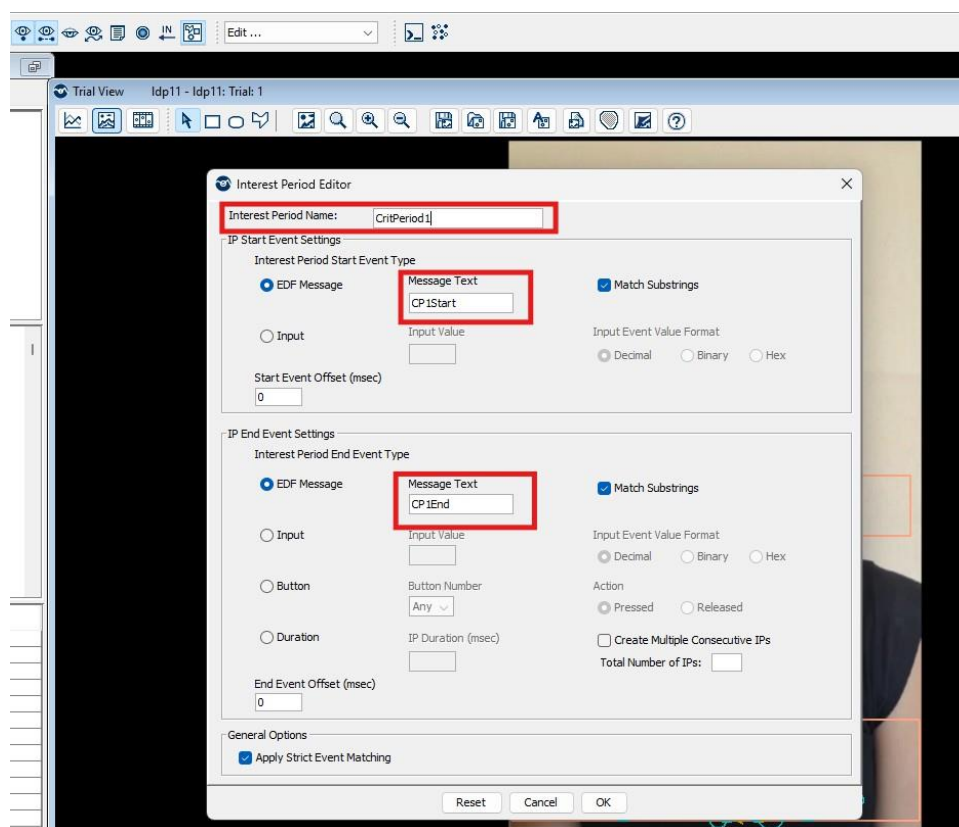
**Figura 32** – Exemplo de comando de importação de mensagens pré-definidas



Fonte: elaborado pelo autor.

Uma vez importados todos os arquivos de mensagens, foram definidos os períodos de interesse no inspetor, a fim de que, em cada lista de mensagens, apenas ficassem visíveis os eventos entre os tempos de entrada e saída de cada período crítico. Para isso, foram usados os eventos exatamente como foram nomeados na lista de mensagem de cada participante.

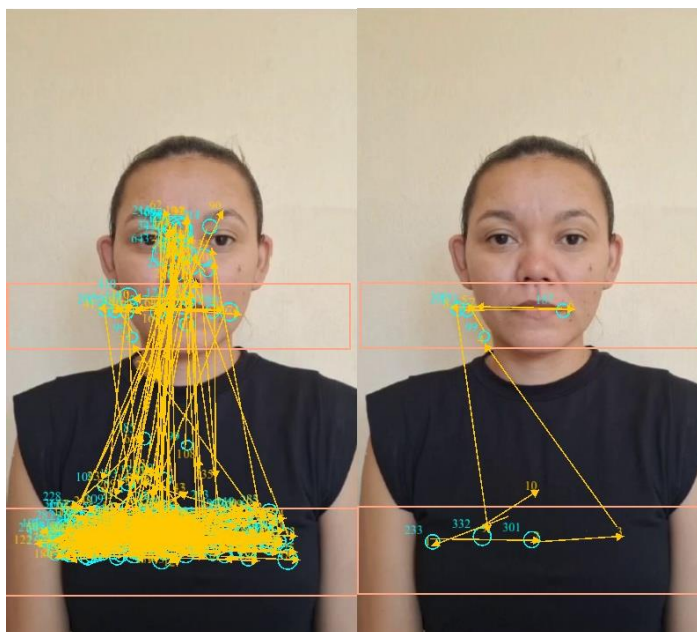
**Figura 33 – Criação de período de interesse no editor**



Fonte: elaborado pelo autor.

Uma vez definidos todos os períodos de interesse, foi possível analisar cada segmento com texto em tela e exportar os dados necessários para tais segmentos. A título de exemplo, pode-se observar a figura abaixo, com dados do participante LDP23. A imagem à esquerda, sem período de interesse definido, apresenta todos os eventos de fixações e sacadas do arquivo EDF; a imagem da direita, no entanto, apresenta apenas os eventos definidos dentro do intervalo do período de interesse 2.

**Figura 34** – Exemplo de mapa de fixações antes e depois



Fonte: elaborado pelo autor.

Como cada vídeo do experimento possui quatro segmentos com texto em tela, foram definidos quatro períodos de interesse. No entanto, durante a limpeza e filtragem dos dados válidos, notou-se que o vídeo 2, na condição parcial, teve um erro de renderização, e o segmento do período de interesse de número quatro também teve a legenda renderizada, fazendo com que o vídeo, nesse segmento, fosse categorizado na condição 1. Sendo assim, optou-se por excluir os dados referentes ao período de interesse 4, a fim de não comprometer os resultados.

A seguir, são apresentadas as medidas de análise utilizadas ao longo da apresentação dos resultados e discussão dos dados, retomadas e adaptadas a partir da seção 2.4.

### **3.8.2 Medidas de Análise**

**Tempo de permanência da área de interesse (IA\_DT):** com base nas condições apresentadas na seção 3.3, retoma-se aqui as suposições feitas na seção 2.4, esperando que:

- Na condição 1, o tempo de permanência relativamente distribuído entre as duas AIs.

- Na condição 2, o tempo de permanência seja maior na AI 2, vista a ausência de legendas na região inferior ao vídeo.
- O tempo de permanência na AI 2 na condição 2 seja menor do que o tempo de permanência na AI 1 na condição 1.

**Duração média de fixações (IA\_MFD) e número total de fixações na área de interesse (IA\_FC):** com base nessas duas médias,

- Na condição 1, legendas totais, espera-se que a duração média das fixações e o número de fixações sejam relativamente distribuídos entre as duas AIs.
- Na condição 2, legendas parciais, espera-se que a duração média de fixações e o número de fixações sejam maiores na AI 2.
- De modo geral, espera-se que a duração média de fixações na AI 2 na condição 2 sejam menores do que os valores obtidos na AI 1 na condição 1, bem como haja mais fixações na AI 2 na condição 2 em comparação à AI na condição 1.

**Tempo da Primeira Fixação na Área de Interesse (IA\_FFD):** para essa medida, espera-se que:

- Tanto na condição 1 quanto na condição 2, a duração média da primeira fixação seja mais longa na AI 2, tendo em vista o surgimento do texto em tela como fator desencadeante de alocação de atenção.
- A duração média da primeira fixação na AI 2 na condição 2 seja menor do que a duração média na AI 2 na condição 1, considerando que a condição 1 gere maior disputa de atenção, justificando uma média maior.

**Análise de Sequência de Fixações:** espera-se que:

- Na condição 1, haja mais transições entre as áreas de interesse AI 1 e AI 2, com menor número de fixações originadas dentro da própria AI;
- Na condição 2, haja menos transições entre as áreas de interesse AI 1 e AI 2, com maior número de fixações originadas da própria AI.

Com base nessas discussões e nas medidas listadas, a seguir, são apresentados os resultados obtidos.

## 4 RESULTADOS

Por se tratar de um estudo piloto, as discussões levantadas neste capítulo são feitas com base em dados preliminares da coleta de dados detalhada acima, da qual participaram seis integrantes. Os dados demográficos dos participantes podem ser conferidos na seção 3.7.3. As discussões são feitas de acordo com cada medida de análise acima mencionada.

No que tange aos valores estatísticos, apesar de o tamanho amostral ser pequeno, o número de testes (vídeos) apresentados justificaram uma análise estatística, com o intuito de observar quão significativos os dados poderiam ser. Para determinar a melhor análise estatística a ser aplicada no presente estudo, utilizou-se o teste de Shapiro-Wilk, que avalia a normalidade da distribuição dos dados. De acordo com Shapiro e Wilk (1965, p. 591), o teste de normalidade avalia a semelhança entre a variância observada nos dados e a variância esperada para uma distribuição normal.

As variáveis quantitativas (i.e., IA\_FFD, IA\_FC, IA\_DT e IA\_MFD) foram testadas quanto à normalidade pelo teste de Shapiro-Wilk. Todas as variáveis apresentaram comportamento não-normal. Portanto, foram realizados testes alternativos não paramétricos (Mann-Whitney-Wilcoxon) para verificar diferenças dessas variáveis respostas em relação à condição (total e parcial) e à área de interesse (*middle* e *bottom*). Os testes foram aplicados separadamente para cada variável categórica. Todas as análises foram realizadas no programa R, adotando-se um nível de significância de  $< 0,05$  para considerar diferenças estatísticas.

### 4.1 Tempo de Permanência na Área de Interesse

O tempo de permanência da área de interesse aqui é medido através do *Dwell Time*, que pode ajudar a entender qual área de interesse chama mais atenção e se há uma preferência visual por uma das AIs em detrimento da outra. Portanto, comparar o tempo de permanência entre as duas legendas pode revelar qual delas tende a ser mais focalizada e por quanto tempo, o que pode ser indicativo de dificuldade de compreensão, interesse ou relevância do conteúdo.

Inicialmente, apresenta-se uma visão geral do padrão de leitura nas duas

áreas de interesse, incluindo ambas as condições, no que tange ao tempo de permanência na área de legenda (área de interesse 1) e no texto em tela (área de interesse 2). Uma limpeza de quatro estágios foi aplicada no tratamento dos dados, para que apenas fossem incluídas fixações entre 80 e 800ms. A Tab. 32, abaixo, resume o tempo médio, incluindo as duas condições.

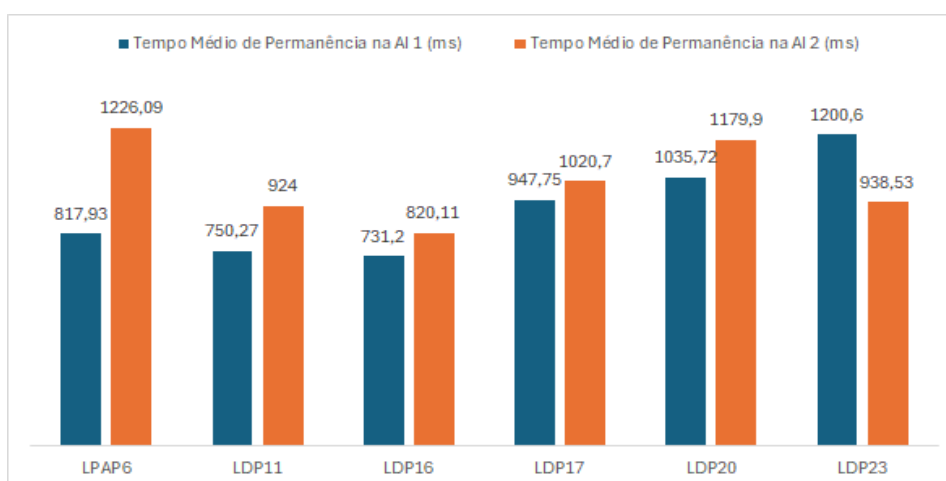
**Tabela 32** – Tempo médio de permanência em cada Área de Interesse em ambas as condições experimentais

| Participante         | Tempo Médio de Permanência (ms) |
|----------------------|---------------------------------|
| AI 1 – Legendas      | 907,33                          |
| AI 2 – Texto em tela | 999,72                          |

Fonte: elaborado pelo autor.

Assim como resume a tabela acima, o tempo médio de permanência na AI 1, considerando ambas as condições experimentais, foi de 907,33 ms, enquanto que na AI 2, essa média foi de 999,72 ms. Ao se observar os valores por participante, mostrados no Gráfico 1, abaixo, nota-se que, com exceção de LDP23, todos os demais participantes seguem na linha dos valores gerais, onde apresentam maior tempo médio de permanência na área do texto em tela, com valores entre 820,11 ms e 126,09 ms. LDP23, no entanto, apresenta um padrão oposto, com média de permanência mais longa na AI 1 (1200,60 ms) em comparação à AI 2 (938,53 ms).

**Gráfico 1** – Tempo Médio de Permanência em cada Área de Interesse por Participante em ambas as condições



Fonte: elaborado pelo autor.

Os valores gerais ajudam a perceber possíveis diferenças em relação ao tempo geral de cada área de interesse, mas os valores divididos por condição podem fornecer *insights* acerca de qual AI pode contribuir para a diferença nos valores, de acordo com cada condição experimental. A Tab. 33, a seguir, resume os tempos médios de permanência em cada AI, de acordo com cada condição de pesquisa.

**Tabela 33 – Tempo Médio de Permanência nas AI por Condição Experimental**

| Condição Experimental | Tempo Médio de Permanência na AI 1 (ms) | Tempo Médio de Permanência na AI 2 (ms) |
|-----------------------|-----------------------------------------|-----------------------------------------|
| Condição 1            | 1127,18                                 | 536,77                                  |
| Condição 2            | 191                                     | 1172,55                                 |

Fonte: elaborado pelo autor.

Na condição 1, o tempo médio de permanência foi de 1127,18 ms na AI 1 e 536,77 ms na AI 2. Na condição 2, os valores apresentaram-se em 191 ms na AI 1 e 1172,55 ms na AI 2. Em outras palavras, na condição 1, houve mais tempo de permanência na AI 1, enquanto na condição 2 ocorreu o contrário, com tempo de permanência mais longo na AI 2.

Essas diferenças também podem ser observadas através dos mapas de calor, de acordo com cada condição experimental. O mapa de calor é criado a partir do mapa de fixação, com base na duração das fixações e contagem das fixações, probabilidade de fixação ao longo da região visível, ou proporção do tempo de permanência do teste (Eyelink, 2024). Sendo assim, os mapas de calor permitem ver as diferenças nas fixações nas áreas de interesse e resumem visualmente a região do vídeo onde há maior e menor densidade de fixações.

Como os períodos de interesse dos vídeos foram definidos de forma não contínua, ou seja, com intervalos de tempo entre um e outro, o *Data Viewer* não consegue exportar um mapa de calor contendo todos os períodos de interesse, mas é possível visualizá-los de forma independente para cada condição experimental. Essa questão foi conversada juntamente com o suporte, que sugeriu à equipe de desenvolvedores essa possibilidade na próxima atualização do programa. De momento, é possível visualizar as diferenças nos mapas de fixações por período de

interesse. Vale ressaltar que apenas os períodos de interesse 1, 2 e 3 estão sendo considerados para a discussão, tendo em vista a exclusão do período de interesse 4, conforme descrito na seção metodológica.

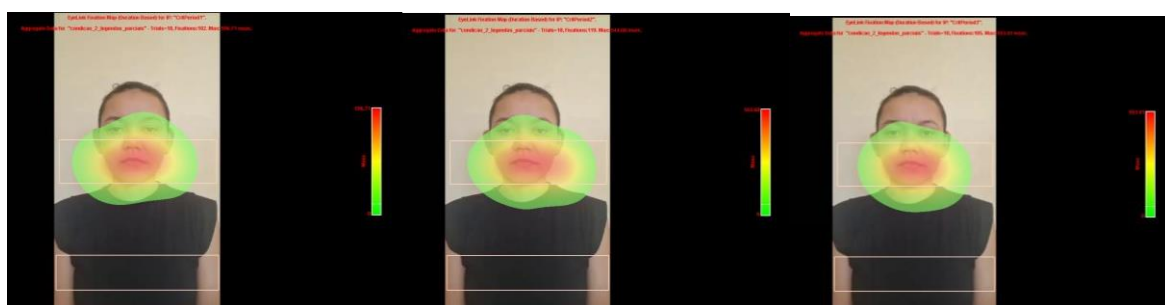
A seguir, traz-se as representações visuais através dos mapas de calor, segundo cada condição experimental. A Fig. 35, a seguir, apresenta a sequência do mapa de fixações através do mapa de calor para a condição 1, e a Fig. 36, a sequência para a condição 2, cujas escalas de cores podem ser lidas de modo que a região em vermelho representa a parte com maior densidade de fixações, e a região em verde, com menor densidade de fixações.

**Figura 35** – Sequência de mapas de calor para a condição 1, “Legendas Totais”, para cada período de interesse



Fonte: elaborado pelo autor.

**Figura 36** - Sequência de mapas de calor para a condição 2, “Legendas Parciais”, para cada período de interesse

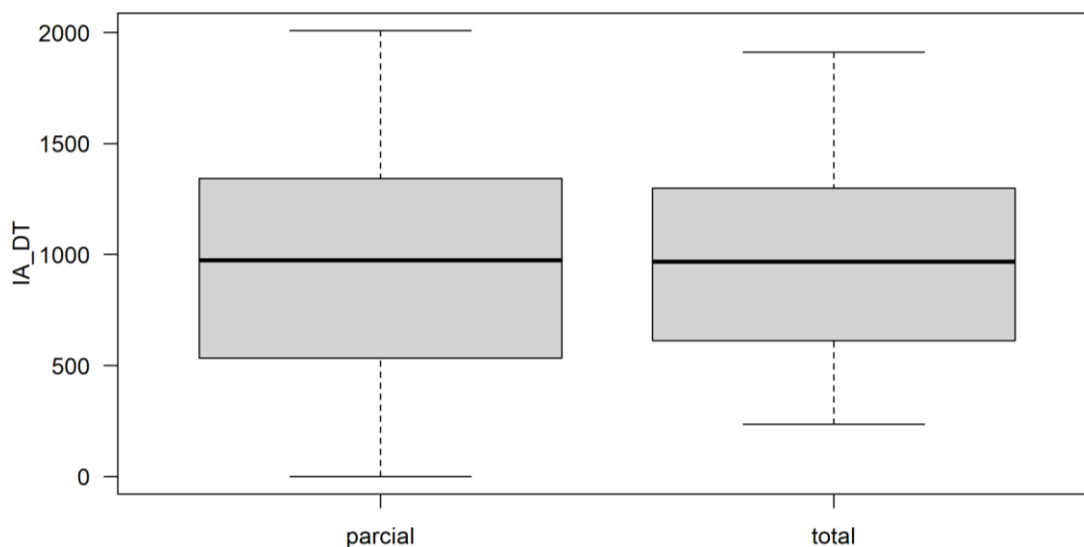


Fonte: elaborado pelo autor.

A partir da Fig. 35, observa-se uma densidade maior de fixações na AI 1, ou seja, na região de legenda, embora com presença de densidade de fixações na AI 2, região do texto em tela, indo ao encontro do que resume a tabela. Na sequência, os mapas de calor para a condição 2 demonstram que os participantes, em sua totalidade, mantiveram uma maior densidade de fixações

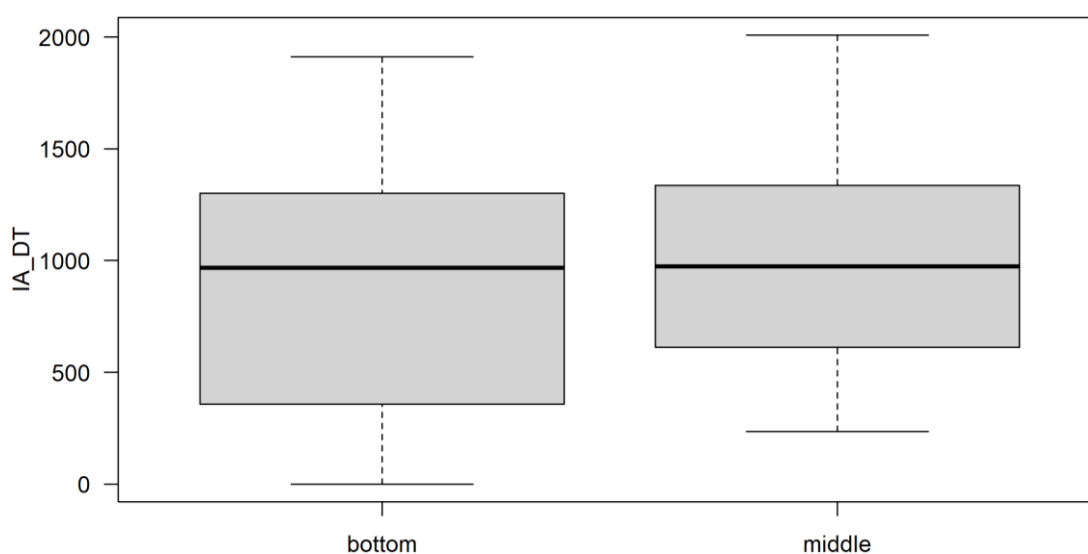
na AI 2, também como resume a Tab. 33. Em relação ao teste Man-Whitney-Wilcoxon, não foram apresentadas diferenças significativas em relação às condições ( $W = 2434$   $p = 0.8350$ ) nem em relação às AIs ( $W = 2236.5$   $p = 0.3085$ ).

**Gráfico 2 – Representação estatística do Tempo Médio de Permanência por Condição Experimental**



Fonte: elaborado pelo autor.

**Gráfico 3 – Representação estatística do Tempo Médio de Permanência por Área de Interesse**

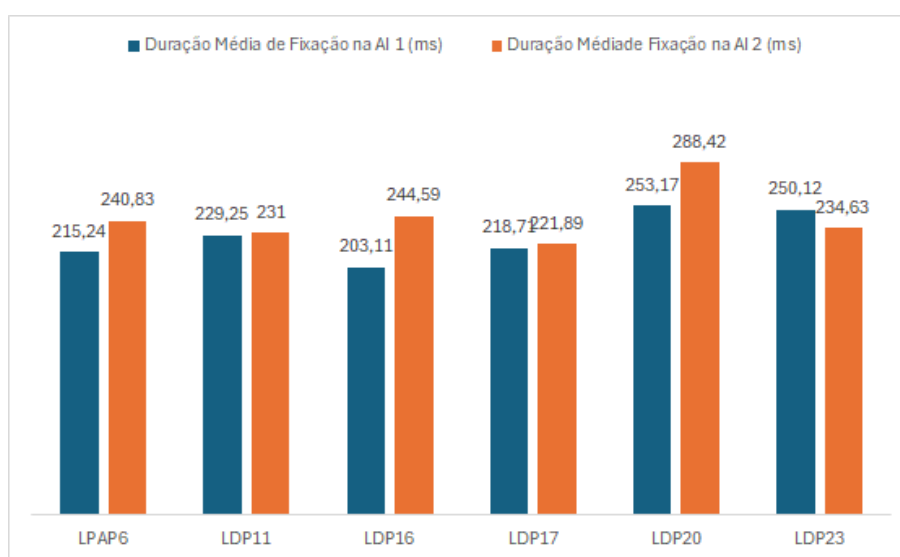


Fonte: elaborado pelo autor.

## 4.2 Médias de Fixações por Área de Interesse

Para essa variável, também se apresentam, inicialmente, os valores gerais, considerando ambas as condições, dando-se seguimento, posteriormente, aos valores obtidos de acordo com cada condição. O Gráfico 4, a seguir, resume os valores médios de fixações por participante para cada área de interesse.

**Gráfico 4 – Duração Média de Fixação por Participante em Ambas as Condições**



Fonte: elaborado pelo autor.

A duração média de fixações na AI 1 variou de 203 ms (LDP16) a 253,17 ms (LDP20); na AI 2, as médias mostraram-se entre 22,89 ms (LDP17) e 288,42 ms (LDP20). Ao se comparar as médias gerais por área de interesse, os valores ficam entre 228,49 ms para a AI 1 e 243,26 ms para a AI 2, como resume a Tab. 34, abaixo.

**Tabela 34 – Duração média de fixações para as Áreas de Interesse AI 1 e AI 2 em ambas as condições experimentais**

| Área de Interesse    | Duração Média de Fixações (ms) |
|----------------------|--------------------------------|
| AI 1 - Legendas      | 228,49                         |
| AI 2 - Texto em Tela | 243,26                         |

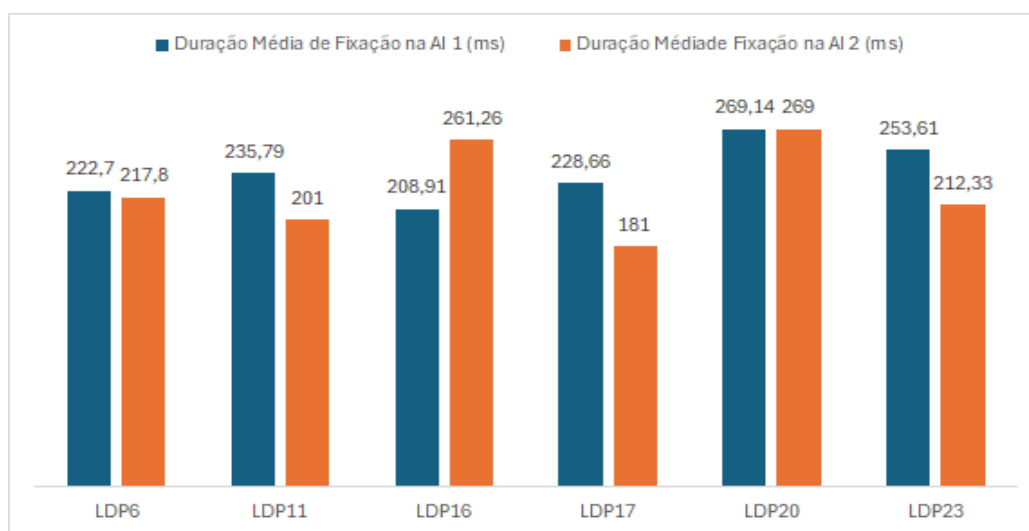
Fonte: elaborado pelo autor.

Além dos valores agrupados por área de interesse, também se realizou cruzamento de dados por condição experimental, com segmentos agrupados por participante, por área de interesse e por período de interesse, considerando não apenas a duração média de fixações, mas também o número médio de fixações.

#### 4.3 Medidas Médias de Fixações por Condição Experimental

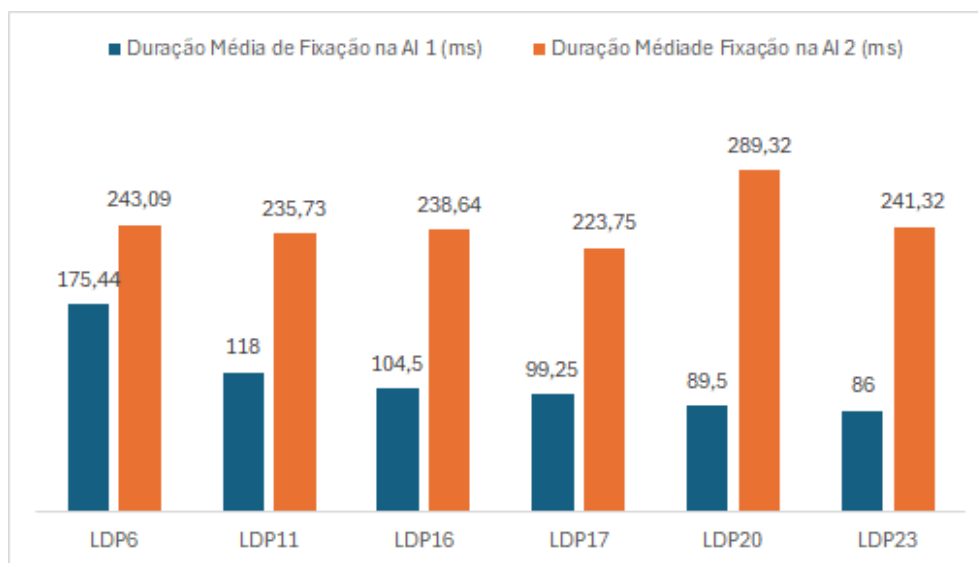
O Gráfico 5, a seguir, apresenta os valores referentes ao primeiro agrupamento, por participante, na condição 1, legendas totais.

**Gráfico 5 – Duração Média das Fixações por participante na condição 1**



Fonte: elaborado pelo autor.

Os valores de duração de fixação obtidos na condição 1 se apresentaram entre 208,91 ms (LDP16) e 269,14 ms (LDP20) na AI 1 e entre 181 ms (LDP17) e 269 ms (LDP20) na AI 2. Na condição 2, legendas parciais, os valores mínimos e máximos ficaram entre 86 ms (AI 1) e 289,32 ms (AI 2). O gráfico abaixo compila as durações médias por participante nesta condição.

**Gráfico 6 – Duração Média das Fixações por participante na condição 2**

Fonte: elaborado pelo autor

Na condição 2, a duração média de fixações ficou entre 86 ms (LDP23) e 175,44 ms (LDP6) na AI 1 e entre 223,75 ms (LDP17) e 289,32 ms (LDP20) na AI 2. Os valores individuais inclinam-se para durações mais longas na AI 2, como se confirma no agrupamento por condição, representado na Tab. 35.

**Tabela 35 – Duração Média de Fixações por AI e por condição experimental**

| Condição Experimental | Duração Média das Fixações na AI 1 (ms) | Duração Média das Fixações na AI 2 (ms) |
|-----------------------|-----------------------------------------|-----------------------------------------|
| Condição 1            | 237,06                                  | 230,04                                  |
| Condição 2            | 130,22                                  | 245,41                                  |

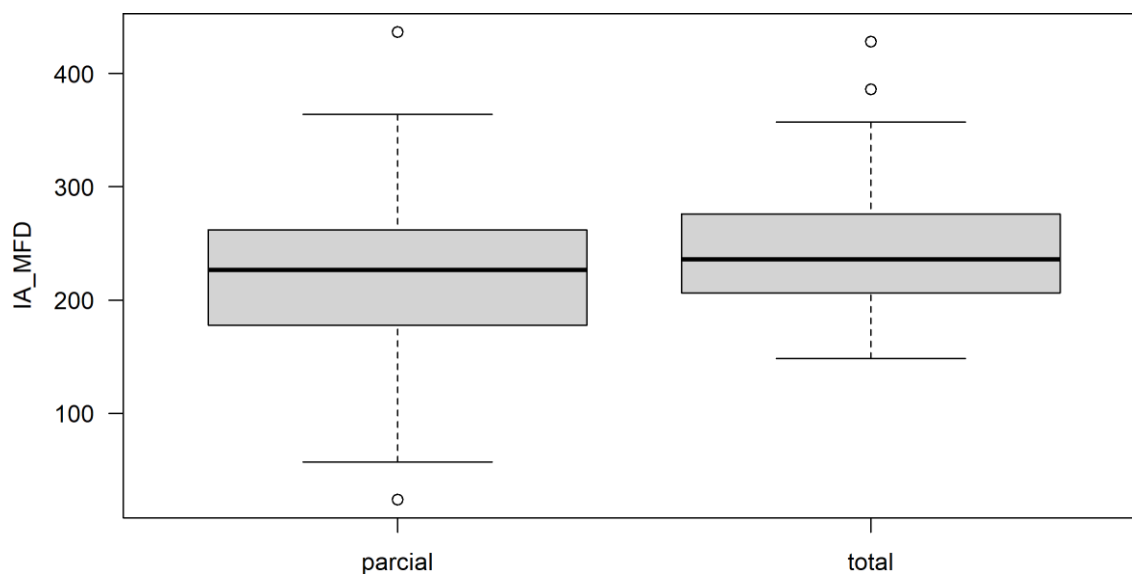
Fonte: elaborado pelo autor.

Na condição 1, a AI 1 apresentou duração média mais longa, com 237,06 ms, em comparação à AI 2, com média de 230,04 ms. Na condição 2, a proporção dos valores se inverte, à medida em que a AI 1 resume média de fixações em 130,22 ms, e a AI 2 apresenta maior duração, com 245,41 ms. Os dados iniciais mostram menor diferença entre os valores da condição 1 em comparação à condição 2. Tais diferenças são apresentadas na seção de discussão.

Do ponto de vista estatístico, não houve diferença significativa em relação

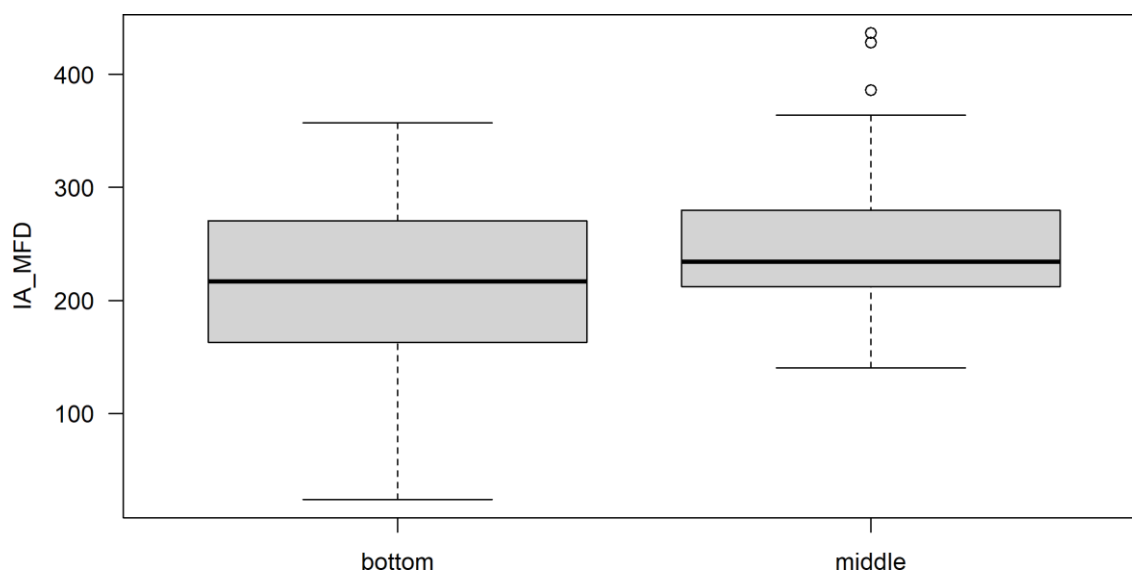
às condições experimentais ( $W = 2090.5$   $p = 0.1351$ ). No entanto, os dados referentes às áreas de interesse ( $W = 1921$   $p = 0.0281$ ) indicaram diferenças significativas para essa variável, como ilustram os gráficos abaixo.

**Gráfico 7 – Dados Estatísticos para a Duração Média de Fixações por Condição Experimental**



Fonte: elaborado pelo autor

**Gráfico 8 – Dados Estatísticos para a Duração Média de Fixações por AI**



Fonte: elaborado pelo autor

As médias de fixações, ou seja, média do número total de fixações na área de interesse, também foram calculadas e são apresentadas na Tab. 36, a seguir.

**Tabela 36** – Média de fixações por AI e por condição experimental

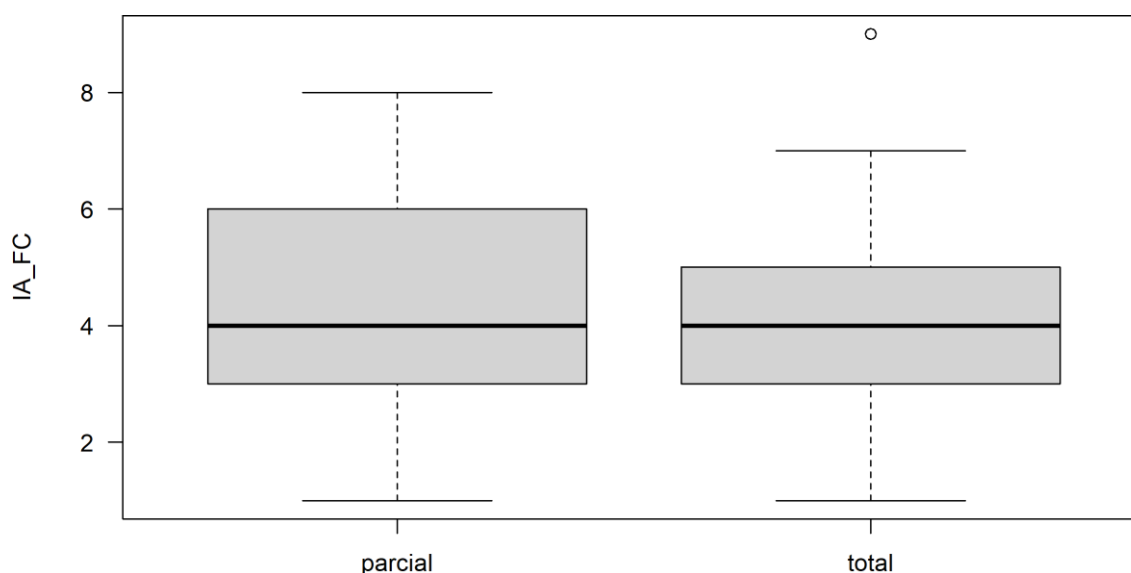
| Condição Experimental | Média de Fixações na AI 1 (ms) | Média de Fixações na AI 2 (ms) |
|-----------------------|--------------------------------|--------------------------------|
| Condição 1            | 42                             | 7                              |
| Condição 2            | 3,6                            | 43                             |

Fonte: elaborado pelo autor.

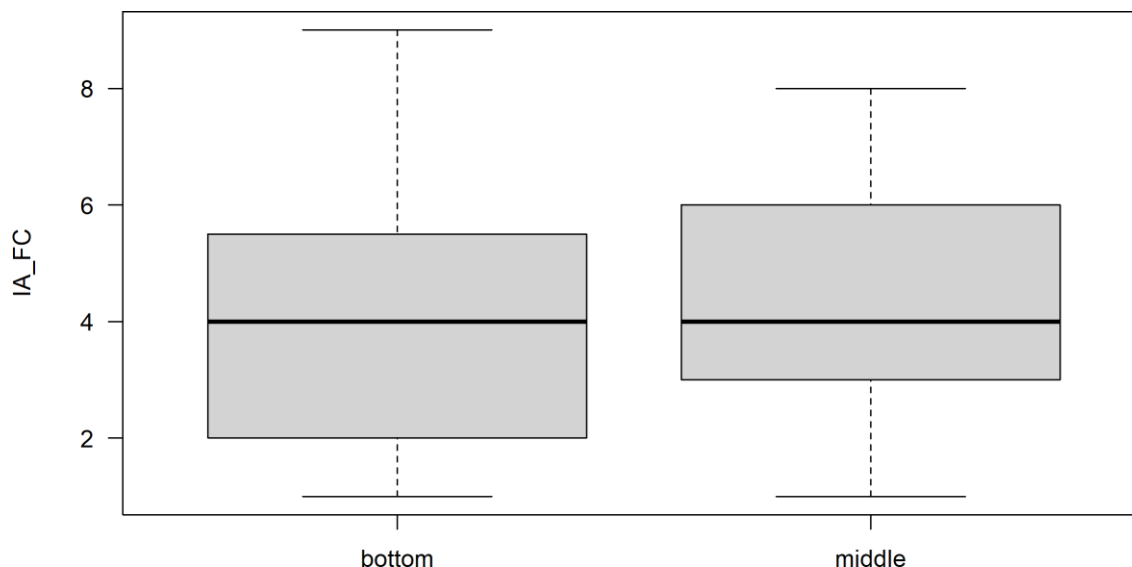
Na condição 1, a AI recebeu uma média total de 42 fixações, enquanto a AI recebeu 7 fixações. Na direção oposta, na condição 2, a AI recebeu 3,6 fixações em média, enquanto a AI 2 recebeu 43.

O teste de Man-Whitney-Wilcoxon não apresentou diferenças significativas entre os valores obtidos por condição experimental ( $W = 5678.5$   $p = 0.7326$ ) nem por área de interesse ( $W = 5531$   $p = 0.5021$ ), como ilustram os gráficos abaixo:

**Gráfico 9** – Dados Estatísticos para a Média de Fixações por Condição Experimental



Fonte: elaborado pelo autor.

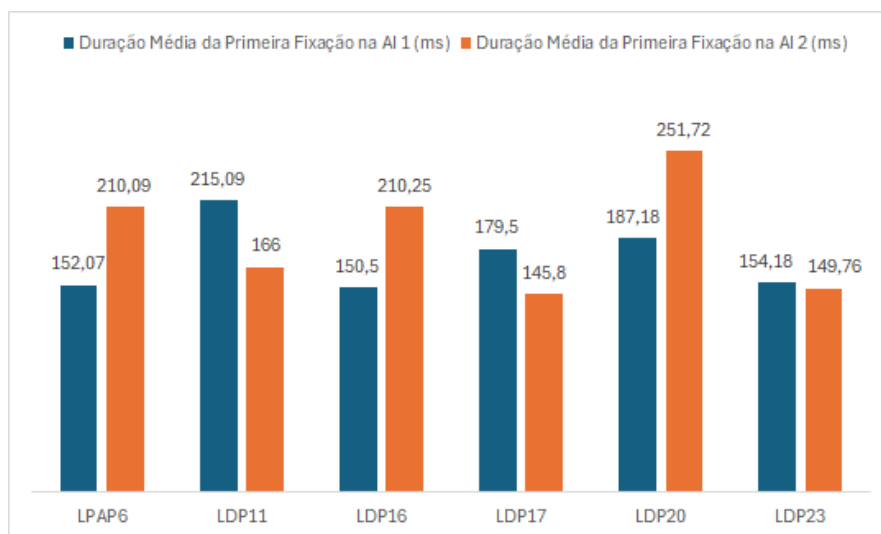
**Gráfico 10 – Média de Fixações por Área de Interesse**

Fonte: elaborado pelo autor.

#### 4.4. Tempo da Primeira Fixação na Área de Interesse

Considerando-se ambas as condições de pesquisa, o Gráfico 11, abaixo, resume os valores médios por participante. As médias da primeira fixação na AI 1 variaram de 150,0 ms (LDP16) a 215,09 ms (LDP11), e, na AI 2, os valores ficaram entre 145,8 ms (LDP20) e 251,72 ms (LDP20), com variações entre os sujeitos. Três participantes (LPAP6, LDP16 e LDP20) apresentaram média de durações mais longas na AI, ao passo que os demais participantes (LDP11, LDP17 e LDP23) tiveram médias de duração mais longas na AI 1.

**Gráfico 11 – Duração média da primeira fixação por participante em ambas as condições**



Fonte: elaborado pelo autor.

Ao se resumir os valores e agrupá-los por área de interesse, observa-se a predominância de primeiras fixações na AI 2, com média de 189,93 ms, em comparação ao resumo de médias na AI 1, calculado em 175,13 ms.

**Tabela 37 – Duração Média da Primeira Fixação por IA em ambas as condições**

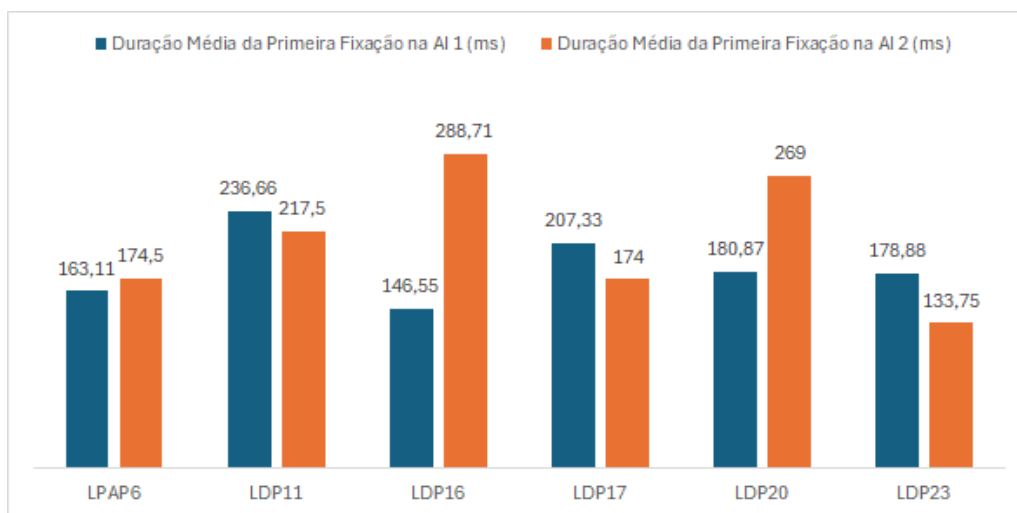
| Condição             | Primeira Fixação (ms) |
|----------------------|-----------------------|
| AI 1 - Legendas      | 175,13                |
| AI 2 - Texto em Tela | 189,93                |

Fonte: elaborado pelo autor.

Embora haja diferenças entre os sujeitos, observa-se uma tendência, ainda que pequena, a fixações iniciais levemente mais longas na região do texto em tela, como resume a Tab. 37 acima, considerando todos os vídeos do experimento e todos os participantes em ambas as condições experimentais. A seguir, são apresentados os valores isolados de cada participante em cada uma das condições experimentais. Ao considerar os dados apenas na condição 1, a média da primeira fixação mostrou valores entre 146,55 ms (LDP16) e 236,66 ms (LDP11) na AI 1. Já na AI 2, os valores médios ficaram entre 133,75 ms (LDP23) e 288,71 ms (LDP16), como

ilustra o gráfico 12, abaixo. Tanto na AI 1 quanto na AI 2, o participante LDP16 apresenta valores limites mínimos e máximos.

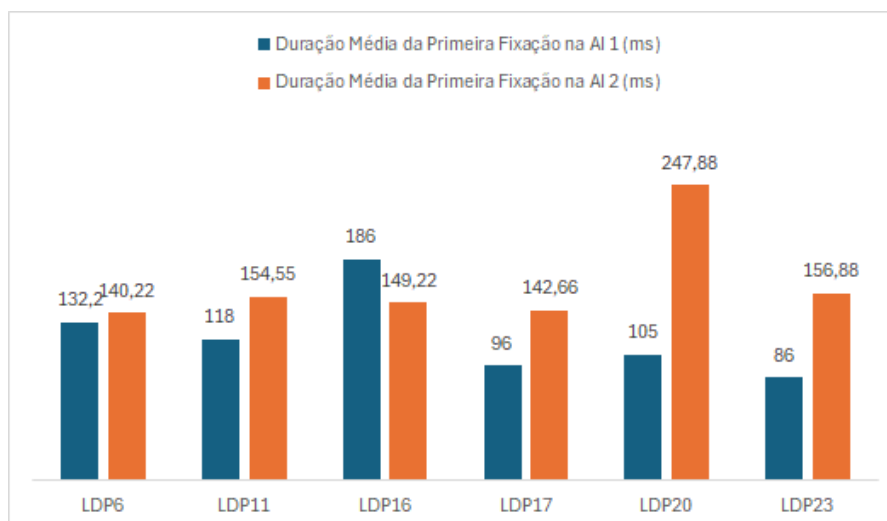
**Gráfico 12 – Duração média da primeira fixação por participante na condição 1**



Fonte: elaborado pelo autor.

Na condição 2, os valores médios na AI 1 mostram-se já relativamente menores do que na condição 1, com médias entre 86 ms (LDP23) e 186 ms (LDP16). Na AI 2, as médias mostram-se entre 140,22 ms (LPAP6) e 247,88 ms (LDP20). Com exceção de LDP16, os demais participantes apresentaram médias mais longas na AI 2.

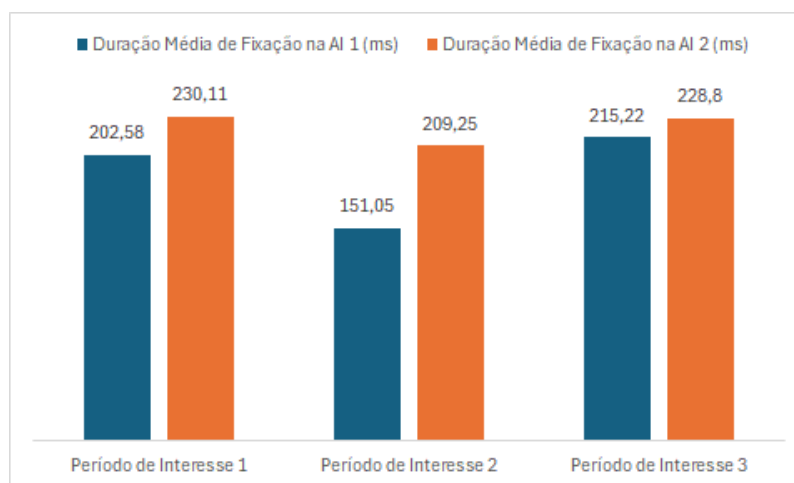
**Gráfico 13 – Duração Média da Primeira Fixação por participante na condição 2**



Fonte: elaborado pelo autor.

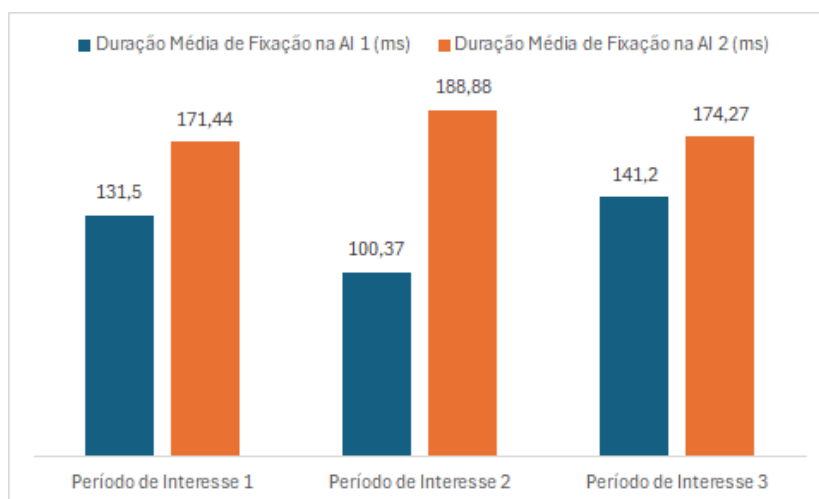
Na sequência, comparam-se também os valores agrupados por período de interesse, a fim de verificar se pode ter havido algum padrão ao longo da exibição dos vídeos, de acordo com cada condição. Os Gráficos 14 e 15, abaixo, resumem as médias da primeira fixação para cada período em cada condição citada.

**Gráfico 14 – Média da Primeira fixação por período de interesse na condição 1**



Fonte: elaborado pelo autor.

**Gráfico 15 – Média da Primeira fixação por período de interesse na condição 2**



Fonte: elaborado pelo autor.

A partir da tabela, observa-se que, na condição 1, os valores dialogam com aqueles apresentados no agrupamento por participantes, apresentando médias de primeira fixação mais longas na AI 2, entre 209,25 ms e 230,11 ms, embora com diferenças pequenas em relação à AI 1, com valores entre 151,05 ms e

215,22 ms. Na condição 2, a duração mais longa na AI 2 ocorreu no período de interesse 2 (188,88 ms) e 3 (174,27 ms) em comparação à AI 1 (100,37 ms e 141,2 ms, respectivamente). Após se agrupar os dados por participante e por período de interesse, agrupa-se também, como principal dado compilatório, os valores por condição experimental. A Tab. 38, abaixo, resume os valores.

**Tabela 38 – Duração média da primeira fixação nas AIs por condição experimental**

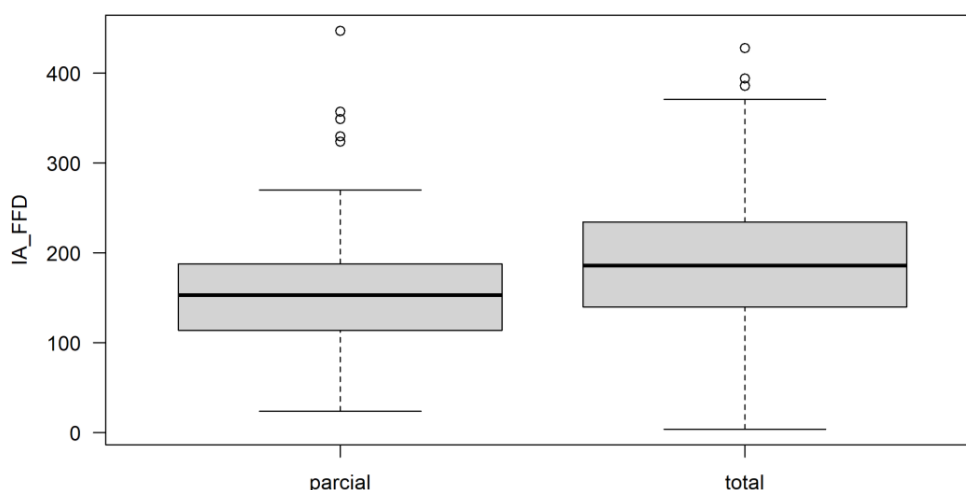
| Condição Experimental | Primeira Fixação na AI 1 (ms) | Primeira Fixação na AI 2 (ms) |
|-----------------------|-------------------------------|-------------------------------|
| Condição 1            | 186,03                        | 225,11                        |
| Condição 2            | 118,13                        | 165,24                        |

Fonte: elaborado pelo autor.

A partir da tabela, observa-se que tanto na condição 1, “legendas totais”, quanto na condição dois, “legendas parciais”, o valor da primeira fixação na AI 2 é mais longo. Na condição 1, a AI 1 teve média de primeira fixação em 186,03 ms, e a AI 2, média de 225,11 ms. Na condição 2, a AI apresentou média de 118,13 ms, enquanto a AI 2 recebeu uma média de 165,24 ms.

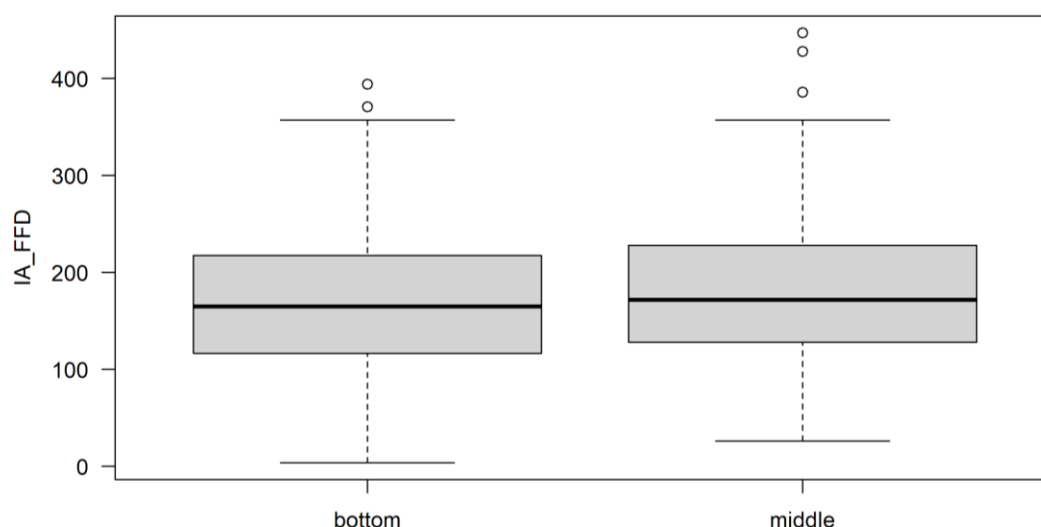
Em relação ao teste de Man-Whitney-Wilcoxon, os valores indicaram diferença significativa em relação às condições ( $W = 1808$   $p = 0.0075$ ), mas não se observou diferenças no que tange às áreas de interesse ( $W = 2235.5$   $p = 0.3767$ ), como ilustram os gráficos abaixo.

**Gráfico 16 – Dados estatísticos para a duração média da primeira fixação nas AIs por condição experimental**



Fonte: elaborado pelo autor.

**Gráfico 17 – Dados estatísticos para a duração média da primeira fixação nas AIs por Área de Interesse**



Fonte: elaborado pelo autor.

Na sequência, são apresentados os resultados referentes à análise da sequência de fixações.

#### 4.5 Análise de Sequência das Fixações

A fim de alinhar os dados obtidos acima para o número médio de fixações e sua duração média, busca-se checar também a sequência de fixações entre as duas áreas de interesse, assim como descrito na seção 3.8.2. Enquanto os valores acima fornecem dados sobre a quantidade de fixações por área de interesse, a sequência de fixações ajuda a entender quantas dessas fixações são provenientes de movimentos de sacadas entre as áreas de interesse.

A Tab. 48, abaixo, traz os valores médios por participante na condição 1, onde:

- IA ID representa a área de interesse analisada;
- IA FSA COUNT 1 representa a média de fixações recebidas na área de interesse AI 1, originadas de todas as áreas de interesse;
- IA FSA COUNT 2 representa a média de fixações recebidas na área de interesse AI 2 originadas de todas as áreas de interesse;
- IA FIX COUNT representa a média de todas as fixações recebidas na área de interesse.

**Tabela 39** – Quantidade média de fixações entre áreas de interesse por participante na condição 1

| IA ID | PARTICIPANTE | IA FSA COUNT 1 | IA FSA COUNT 2 | IA FIX COUNT |
|-------|--------------|----------------|----------------|--------------|
| 1     | LPAP6        | 3,88           | 1              | 5,33         |
| 2     | LPAP6        | 1              | 1,5            | 2,5          |
| 1     | LDP11        | 3              | 1              | 3,77         |
| 2     | LDP11        | 1              | 2              | 3            |
| 1     | LDP16        | 2,37           | 1              | 3,77         |
| 2     | LDP16        | 1              | 1,5            | 2,14         |
| 1     | LDP17        | 4,33           | 1              | 5,33         |
| 2     | LDP17        | 0              | 1              | 2            |
| 1     | LDP20        | 3,75           | 0              | 5,12         |
| 2     | LDP20        | 1              | 0              | 1            |
| 1     | LDP23        | 3,88           | 1              | 5,22         |
| 2     | LDP23        | 1              | 2              | 3            |

Fonte: elaborado pelo autor.

Na condição 1, a média total de fixações na AI 1 apresentou números entre 3,77 (LDP11 e LDP16) e 5,33 (LPAP6 e LDP17). Na AI 2, a média total de fixações ficou entre 1 fixação (LDP20) e 3 fixações (LDP11 e LDP23). Desse total médio, o participante LDP17, na AI 1, foi responsável pelo maior número de fixações sequenciadas dentro da própria área de interesse, ou seja, das 5,33 fixações médias recebidas na AI 1, 4,33 foram feitas com sacadas dentro na própria área, com 0 fixações originadas com sacadas vindas da AI 2. E, das 2 fixações médias na AI 2, 1 fixação foi feita após uma sacada dentro da própria área de interesse, e uma fixação após sacadas vindas AI 1. O participante com menor número de fixações sequenciadas na própria área de interesse foi LDP16, que, das 3,77 fixações na AI 1, 2,37 foram feitas com sacadas dentro da própria AI. E, das 2,14 fixações na AI 2, 1,5 fixações foram feitas com sacadas dentro da própria AI, enquanto 1 fixação foi feita após sacada vinda da AI 1. Na sequência, a Tab. 40 traz os valores médios por participante na condição 2.

**Tabela 40** – Quantidade média de fixações entre áreas de interesse por participante na condição 2

| IA ID | PARTICIPANTE | IA FSA COUNT 1 | IA FSA COUNT 2 | IA FIX COUNT |
|-------|--------------|----------------|----------------|--------------|
| 1     | LPAP6        | 1,5            | 1              | 1,8          |
| 2     | LPAP6        | 1              | 5,25           | 5,66         |
| 1     | LDP11        | 0              | 1              | 1            |
| 2     | LDP11        | 0              | 3,22           | 4,22         |
| 1     | LDP16        | 1              | 0              | 2            |
| 2     | LDP16        | 0              | 3,44           | 4,66         |
| 1     | LDP17        | 1              | 1              | 1,33         |
| 2     | LDP17        | 1              | 3,77           | 4,88         |
| 1     | LDP20        | 1              | 0              | 1,33         |
| 2     | LDP20        | 1              | 3,77           | 4,77         |
| 1     | LDP23        | 0              | 1              | 1            |
| 2     | LDP23        | 0              | 3,44           | 4,44         |

Fonte: elaborado pelo autor.

Na condição 2, a média total de fixações na AI 1 apresentou números entre 0 (LDP16) e 1,8 fixações (LPAP6). Na AI 2, a média total de fixações ficou entre 4,22 (LDP11) e 5,66 fixações (LPAP6). Desse total médio, a AI 2 recebeu o maior número de fixações sequenciadas com sacadas dentro da própria área de interesse, com os dados de LPAP6, que, do total médio de 5,66 fixações, 5,25 foram feitas com sacadas dentro da própria AI, e uma fixação foi feita a partir da AI 1. E, das 1,8 fixações recebidas na AI 1, 1,5 foram feitas a partir de sacadas dentro da própria AI, com média de uma fixação vinda da AI 2. No caso contrário, o participante com menor número de fixações na AI 2 foi LDP11 que, das 4,22 fixações gerais recebidas na área de interesse, 3,22 foram feitas com sacadas dentro da própria AI e 1 fixação partindo da AI 1. Na AI 1, o participante apresentou média de apenas uma fixação feita na própria área de interesse, sem fixações vindas da AI 2. Os valores médios também são compilados nas Tab. 41, abaixo.

**Tabela 41** – Quantidade média de fixações entre áreas de interesse por área de interesse na condição 1

| Área de Interesse | IA_FSA_COUNT_1 | IA_FSA_COUNT_2 | IA FIX COUNT |
|-------------------|----------------|----------------|--------------|
| AI 1              | 3,56           | 1              | 4,75         |
| AI 2              | 1              | 1,69           | 2,33         |

Fonte: elaborado pelo autor.

A partir do número médio de fixações nas áreas de interesse AI 1 e AI 2 na condição 1, observa-se que, em média, a cada 4,75 fixações realizadas em AI 1, 3,56 fixações partiram da própria área de interesse AI 1, e apenas 1 fixação foi originada a partir da AI 2. Em relação à AI 2, a cada 2,33 fixações recebidas, 1,69 fixações se originaram na própria AI 2, e uma fixação foi feita a partir da AI 1, indicando que houve maior sequência de fixações dentro da AI 1 em comparação à AI 2, o que pode sugerir menor engajamento de leitura contínua na AI 2.

**Tabela 42** – Quantidade média de fixações entre áreas de interesse por área de interesse na condição 2

| Área de Interesse | IA_FSA_COUNT_1 | IA_FSA_COUNT_2 | IA FIX COUNT |
|-------------------|----------------|----------------|--------------|
| AI 1              | 1,2            | 1              | 1,46         |
| AI 2              | 1              | 3,72           | 4,77         |

Fonte: elaborado pelo autor.

Na condição 2, na área de interesse AI 1, a cada 1,46 fixações recebidas, 1,2 fixações foram realizadas dentro da própria área de interesse e 1 fixação foi originada a partir da AI 2. Na área de interesse AI 2, a cada 4,77 fixações recebidas, 3,72 fixações ocorreram dentro da mesma área de interesse, enquanto apenas 1 fixação ocorreu a partir da AI 1. A seguir, apresentam-se os resultados obtidos no Questionário pós-coleta.

## 4.6 Questionário Pós-Coleta

Como especificado na seção metodológica, após a coleta de dados, cada participante respondeu ao questionário Pós-Coleta (ver Apêndice B). O questionário consistiu de 15 perguntas com respostas entre “verdadeiro” ou “falso”, a fim de se avaliar o grau de compreensão dos sujeitos sobre o assunto dos vídeos apresentados. Sendo assim, cada resposta correta seria computada como “1”, e cada resposta incorreta seria computada como “0”, como ilustra a figura abaixo. Ao final de todas as 15 perguntas, cada participante receberia uma pontuação entre 0 e 15, de acordo com a quantidade de acertos.

**Figura 37 – Sistema de pontuação por sujeito para o questionário pós-coleta**

Color

≡

Share and sync

| <div> <div>○</div> <div>Afirmativa 3 - Defini...</div> <div>▼</div> </div> | <div> <div>f₃</div> <div>Afirmativa 3</div> <div>▼</div> </div> | <div> <div>f₄</div> <div>Pontuação</div> <div>▼</div> </div> | <div> <div>☺</div> <div>Afirmativa 4 - A for...</div> <div>▼</div> </div> |
|----------------------------------------------------------------------------|-----------------------------------------------------------------|--------------------------------------------------------------|---------------------------------------------------------------------------|
| Falso                                                                      | 1                                                               | 5                                                            | Falso                                                                     |
| Verdadeiro                                                                 | 0                                                               | 3                                                            | Falso                                                                     |
| Falso                                                                      | 1                                                               | 5                                                            | Falso                                                                     |
| Falso                                                                      | 1                                                               | 5                                                            | Falso                                                                     |
| Falso                                                                      | 1                                                               | 5                                                            | Falso                                                                     |
| Falso                                                                      | 1                                                               | 5                                                            | Falso                                                                     |
| Falso                                                                      | 1                                                               | 5                                                            | Falso                                                                     |
| Falso                                                                      | 1                                                               | 4                                                            | Falso                                                                     |
| Falso                                                                      | 1                                                               | 5                                                            | Falso                                                                     |
| Falso                                                                      | 1                                                               | 4                                                            | Falso                                                                     |
| Falso                                                                      | 1                                                               | 5                                                            | Falso                                                                     |

Fonte: elaborado pelo autor.

As pontuações de cada participante são compiladas na tabela abaixo:

**Tabela 43 – Total de acertos por participante**

| Participante | Número de Acertos | % de Acertos |
|--------------|-------------------|--------------|
| LPAP6        | 15                | 100.00%      |
| LDP11        | 15                | 100.00%      |
| LDP16        | 13                | 86.67%       |
| LDP17        | 15                | 100.00%      |
| LDP20        | 13                | 86.67%       |
| LDP23        | 15                | 100.00%      |

Fonte: elaborado pelo autor.

Como pode ser visto na tabela, quatro dos participantes acertaram todas as perguntas do questionário, e dois participantes (LDP16 e LDP20) acertaram 13 das 15 perguntas. Ao se observar de forma individual as respostas incorretas, tem-se as afirmativas 5 e 11 (A5 e A11) para LDP16 e as afirmativas 10 e 11 (A10 e A11) para LDP20.

A afirmativa A5 está ligada às informações presentes no vídeo 2, e as afirmativas A10 e A11 estão ligadas às informações presentes no vídeo 4. Em relação à condição experimental em que os participantes assistiram tais vídeos, LDP16 os viu na condição 2, “legendas parciais”, para A5, e na condição 1, “legendas totais”, para A11; já para LDP20, ambas as afirmativas foram vistas na condição 1. No que diz respeito ao momento do vídeo em que as respostas às afirmativas poderiam ser recuperadas, A5 poderia ser recuperada do vídeo 2 no período de interesse CP3; A10 poderia ser recuperada do vídeo 4, em CP 1; A11, a partir do vídeo 4, no CP3. O quadro 11, abaixo, resume como cada afirmativa foi distribuída para ambos os participantes, de acordo com cada condição experimental e período de interesse nos vídeos.

**Quadro 11 – Afirmativas erradas por participante e vídeo relacionado**

| Participante | Afirmativa | Vídeo   | Condição        | Período de Interesse |
|--------------|------------|---------|-----------------|----------------------|
| LDP16        | A5         | Vídeo 2 | 2 - L. Parciais | CP3                  |
|              | A11        | Vídeo 4 | 2 - L. Totais   | CP3                  |
| LDP20        | A10        | Vídeo 4 | 2 - L. Totais   | CP1                  |
|              | A11        | Vídeo 4 | 2 - L. Totais   | CP3                  |

Fonte: elaborado pelo autor.

## 5 DISCUSSÕES INICIAIS

Neste capítulo, será feita a discussão dos dados apresentados na seção anterior, de acordo com as medidas de análise previamente apresentadas.

### 5.1 Tempo de permanência na área de interesse

De modo geral, os valores médios de permanência apresentam-se relativamente próximos para LDP11, LDP16, LDP17 e LDP20. Observa-se que LPAP6 e LDP23, no entanto, possuem valores maiores de diferença entre a AI 1 e a AI 2. A descrição qualitativa do comportamento ocular dos sujeitos pode fornecer dados interessantes na avaliação da metodologia empregada nesse tipo de condição experimental. E, nesse caso, os dados gerais iniciais mostram diferenças entre os sujeitos participantes da pesquisa.

Além disso, os valores gerais ajudam a perceber possíveis diferenças em relação ao tempo geral de cada área de interesse, mas os valores divididos por condição podem fornecer *insights* acerca de qual AI em que condição pode contribuir para a diferença nos valores.

Retomando-se as hipóteses apresentadas anteriormente, para essa medida, era esperado que, na condição 1, houvesse tempo de permanência relativamente distribuído entre as duas AIs, tendo em vista a possibilidade de distribuição de atenção entre dois focos de informação redundante apresentados de forma simultânea; na condição 2, esperava-se que houvesse tempo de permanência maior na AI 2, considerando a inexistência de texto na AI 1.

Retomando os valores resumidos por condição, observou-se, como esperado, que o tempo médio de permanência na AI 1 é maior na condição 1 (1127,18 ms), enquanto na condição 2 é menor (536,77 ms). Essas diferenças também puderam ser observadas através dos mapas de calor de acordo com cada condição experimental. Assim como já previsto, a região da legenda, na AI 1, teve uma maior densidade de fixações em todos os períodos de interesse. Dando prosseguimento, em relação à condição 2, esperava-se o oposto, tendo em vista a não existência das legendas nesses segmentos, bem como os menores tempos de permanência expostos na tabela.

Embora os números gerais estejam dentro do esperado, nota-se o seguinte:

- 5.1.1 Observação 1 – ocorrência de atividade ocular na AI 1 na condição 2: mesmo não havendo legendas na AI 1 na condição 2, há atividade ocular nessa área. Por ser um valor de média baixo, principalmente em comparação à AI 2, sugere-se que sejam devido a fixações iniciais na área de legenda, a nível de varredura rápida, uma vez que o participante já venha com o fluxo de leitura nesta região, buscando pela continuidade da legenda ou informação nova, e, então, direciona o olhar para a AI 2, onde há informação textual, com fixações mais longas.
- 5.1.2 Observação 2 – Maior tempo de permanência na AI 1 na condição 1: um valor maior de *Dwell Time* na AI, claramente, indica que os participantes passaram mais tempo na região da legenda, o que pode se dar por duas razões. Uma primeira explicação seria que o texto presente na AI poderia conter mais informações, precisando de mais tempo para ser processado. No entanto, levando em consideração a seleção dos textos dos roteiros e o controle de palavras e velocidade das legendas, além do fato de que os textos presentes em ambas as AIs serem idênticos na condição 1, como descrito na seção metodológica, descarta-se inicialmente, essa hipótese de análise. Outra explicação poderia ser percebida em decorrência de a região das legendas serem percebidas pelos participantes como mais importante ou relevante em comparação à AI 2 e, portanto, passarem mais tempo nessa região do vídeo, dialogando com os resultados descritos por Liao, Kruger e Doherty (2020), ao compararem a leitura de legendas bilíngues. Uma análise do número de fixações a cada AI se torna relevante nesse sentido e é apresentada mais adiante para corroborar ou não essa hipótese de análise.
- 5.1.3 Observação 3 – maior tempo de permanência na AI 2 na condição 2 em comparação ao tempo de permanência na AI na condição 1: levando-se em consideração a observação 2, acima, a legenda poderia ser entendida como fonte principal de informação na condição 1. Contudo, ao ser retirada na condição 2, dando espaço apenas ao texto em tela, este passa a ser a fonte principal de informação no vídeo. Dessa forma, é justo que se compare esses tempos de permanência de AI 1 na condição 1 e

em relação ao tempo na AI 2 na condição 2. O que se observa, inicialmente, é que, embora, sejam próximos, há uma pequena diferença, com maior tempo na AI 2 na segunda condição, o que pode indicar maiores tempos de fixação nesta região. Para tanto, discorre-se sobre esses valores na seção 5.1.

Estes resultados também podem ser analisados à luz do modelo de leitura automática descrito por d'Ydewalle e De Bruycker (2007), que argumentam que a presença de legendas atrai o olhar automaticamente, estabelecendo-as como a principal fonte de informação visual. Este comportamento justifica o maior tempo de permanência na AI 1 na condição de legendas totais, mesmo na presença de textos redundantes em tela. Por outro lado, o modelo de Kruger, Hefer e Matthew (2013), que propõe a redução do esforço cognitivo por meio de layouts otimizados para evitar redundâncias visuais, ajuda a explicar o maior tempo de permanência na AI 2 na condição parcial, onde a ausência de legendas promove um padrão de fixações mais fluido na área de texto em tela.

Sendo assim, esses dados iniciais sugerem que, de modo genérico, alguns participantes passaram mais tempo na região de legenda, enquanto outros o fizeram na região do texto em tela. No entanto, faz-se necessário buscar os valores médios de duração de fixações, como apresentado na seção seguinte, por cada AI e também por cada condição experimental.

## **5.2 Medidas Médias de Fixações**

A duração média de fixações é uma medida comumente analisada em pesquisas com rastreamento ocular, uma vez que está diretamente ligada ao princípio da ligação olho-mente (JUST; CARPENTER, 1980, apud ASSIS, 2021, p. 121), sugerindo que um maior tempo de fixação possa estar relacionado a um maior custo de processamento. Retomando as hipóteses formuladas para essa medida de análise, esperava-se que:

- Na condição 1, legendas totais, a duração média das fixações e o número de fixações seja relativamente distribuído entre as duas AIs, em decorrência da divisão de atenção esperada entre as duas áreas de interesse.
- Na condição 2, legendas parciais, a duração média de fixações e o número

de fixações sejam maiores na AI 2, levando em consideração a ausência de texto na AI 1.

- A duração média de fixações na AI 2 na condição 2 sejam menores do que os valores obtidos na AI 1 na condição 1.

Esse comportamento ocular reflete as hipóteses dos modelos de leitura descritos por d'Ydewalle e De Bruycker (2007), que destacam a predominância de legendas como foco principal, especialmente em condições de competição visual com outros elementos. Na condição de legendas totais, o aumento nas durações médias na AI 1 pode ser associado à necessidade de integrar informações textuais e visuais em um curto espaço de tempo. Em contrapartida, na condição de legendas parciais, o padrão de fixações mais uniforme na AI 2 dialoga com o modelo de Kruger, Hefer e Matthew (2013), que enfatiza o papel do design para minimizar custos cognitivos, otimizando a leitura em contextos audiovisuais.

Os resultados gerais da média de fixações em cada uma das Ais, independentemente da condição de pesquisa, podem fornecer informações preliminares sobre a movimentação ocular dos sujeitos em relação à região do vídeo. De modo geral, a partir dos valores individuais, os participantes apresentaram durações mais curtas na AI 1, com valores entre 203 ms, para LDP16, e 253,17 ms, para LDP20. O participante LDP23, no entanto, apresenta valores diferentes dos demais, com média de fixação maior na AI 1 (250,12 ms), embora com diferenças pequenas entre as duas Ais. Além disso, nota-se que a diferença entre as médias em ambas as condições são relativamente pequenas de forma geral.

Em suma, os dados individuais dialogam com o panorama apresentado por área de interesse, com exceção de LDP23, como pontuado. Na seção anterior, viu-se que o tempo médio de permanência na AI 2 (999,72 ms), considerando os valores em ambas as condições experimentais, foi maior em comparação ao tempo de permanência na AI 1 (907,33 ms). A partir dessa lógica, era esperado que a duração média das fixações em AI 1 seguisse na mesma direção, tendo em vista que são variáveis relacionadas. Com base nos valores da Tab. 35, inicialmente, o panorama geral para as duas áreas de interesse, levando-se em consideração as duas condições experimentais, mostrou tempos de fixação mais

longos na AI 2 (243,26 ms) em comparação à AI 1 (228,49 ms). Sendo assim, os valores aparecem dentro do esperado.

Embora ambas as condições experimentais sejam consideradas aqui, observa-se que há certa diferença entre os dois valores, assim como previamente exposto na seção 4.1, com uma predominância de fixações mais longas na AI 2, mas com diferenças pequenas, considerando que a AI 1 na condição 2 não continha informação textual.

Com base nesse valor, pode-se fazer a seguinte observação, a fim de se entender o porquê dessa diferença. Assim como descrito na literatura, fixações mais longas podem indicar maior custo de processamento, o que pode estar diretamente ligado (não limitante) ao nível de complexidade da informação apresentada. Sendo assim, a primeira impressão sobre o porquê de a média de fixações ter sido maior na AI 2, região do texto em tela, poderia sugerir uma maior complexidade textual. No entanto, pelo mesmo motivo acima descrito, descarta-se também essa inferência, uma vez que os textos são iguais em ambas as AIs na condição 1. Em segundo lugar, pode haver maior dificuldade no processamento das informações presentes na AI 2, resultando em fixações mais longas.

Como as médias de permanência se mostraram mais longas na condição 2, é pertinente investigar qual condição apresenta maior média de fixação por participante.

### ***5.2.1 Médias de Fixações por Condição Experimental***

Retomando o que era esperado dentro dessa medida de análise, a duração média de fixações para a AI 1 foi de 237,06 ms, e, para a AI 2, 230,04 ms. Na condição 2, a AI 1 teve uma média de 130,22 ms, enquanto a AI 2 teve 245,41 ms, ou seja, dentro do esperado para ambas as condições. O mesmo se reflete em relação ao número total médio de fixações por AI, onde a AI 1 recebeu mais fixações na condição 1 e menos fixações na condição 2.

No entanto, observa-se que, na duração média de fixações, os valores referentes à AI 2 em ambas as condições e em ambas as variáveis se apresentam relativamente próximos. Por outro lado, o número de fixações, nesse sentido, apresenta valores com diferenças maiores, como na AI 2, na condição 1, por exemplo, embora tenha recebido poucas fixações ao longo do vídeo,

apresenta uma duração média de fixações relativamente próxima à média AI 2 da condição 2, que recebeu mais visitas e, portanto, possui um maior tempo de distribuição de fixações.

Ao se observar os valores individuais, nota-se que eles vão ao encontro do que se esperava nesta condição em alguns casos, com duração média de fixações relativamente distribuídas entre as duas AIs, principalmente no tange aos participantes LPAP6 e LDP20, cujos valores são bastante próximos. No caso de LDP20, a diferença é quase inexistente, 269,14 ms na AI 1 e 269 ms na AI 2. LDP16, no entanto, apresenta valores na direção oposta, sendo o único participante que mostrou duração média de fixações maiores na AI 2 nesta condição.

**Tabela 44 – Medidas de Fixação de LDP16 na condição 1**

| Medidas               | AI 1   | AI 2   |
|-----------------------|--------|--------|
| Média de Fixação (ms) | 208,91 | 261,26 |
| Número de fixações    | 34     | 15     |

Fonte: elaborado pelo autor

Ao se analisar os números de fixações do participante, nota-se que ele teve menos visitas em termos de fixações na AI 2, ocasionando o aumento das durações médias. Retomando o que se resumiu na seção 3.7.3, durante o preenchimento do questionário pré-coleta, na pergunta “*Você tem alguma dificuldade em ver filmes/programas legendados?*”, LDP16 respondeu que sim e, na justificativa, “*Por favor especifique qual dificuldade você tem ao ver filmes legendados*”, ele afirmou ter dificuldade em se concentrar ao assistir esse tipo de material. Apesar de os dados serem insuficientes para tirar conclusões, infere-se que a duração média de fixações mais longas na AI 2 possa ter sido ocasionada por dificuldade de concentração na AI 1, como fonte principal de informação, como foi o caso dos demais sujeitos. Dessa forma, ao haver o surgimento do texto em tela na AI 2, LDP16 teve sua atenção direcionada a esse ponto do vídeo. Tal situação também pode ser justificada com base no que afirma Jensema et al (2000), em que o aparecimento de legendas leva o espectador a olhar diretamente a elas, sem muito controle sobre o ponto de atenção. Neste caso,

a situação é causada pelo texto em tela.

De forma semelhante, como esperado, as durações médias na AI 2 na condição parcial foram maiores, uma vez que as legendas da AI foram retiradas, levando os participantes a focarem na AI 2 como fonte de informação predominante. O teste estatístico de Mann-Whitney-Wilcoxon, aplicado para avaliar essa diferença, apresentou um valor significativo (**W = 1921, p = 0.0281**), o que reforça que essa diferença não ocorreu ao acaso. O intervalo de confiança para o rank (-0.22; 95%CI [-0.39, -0.03]) confirma que a posição da legenda exerce um impacto mensurável sobre a duração média das fixações.

Os resultados sugerem que a posição da legenda influencia a interação visual do espectador, com maior atenção sendo dedicada à legenda posicionada no centro da tela. Essa maior atenção pode estar relacionada ao equilíbrio entre o processamento da legenda e os elementos visuais centrais do vídeo.

O foco maior em AI 1 também pode indicar que os participantes acham as informações nessa área mais fáceis de processar quando ambas as fontes de texto estão disponíveis. Nesse caso em particular, o maior tempo de fixação não deve ser considerado como um dado de maior custo de processamento, mas como fonte principal de informação, o que pode indicar que, mesmo com outra fonte de informação em tela, os participantes parecem optar pela área de legenda como fonte de informação principal. Sendo assim, embora as condições experimentais sejam diferentes, esses resultados se assemelham aos dados obtidos por Liao, Kruger e Doherty (2020).

### ***5.2.2 Médias de fixação por período de interesse***

Além dos valores médios de fixação por sujeito, cogitou-se também se haveria padrões de fixações à medida que o vídeo fosse assistido, ou seja, como os participantes teriam sua atenção distribuída em ambas as condições por período de interesse. Seria possível perceber o padrão do design de legenda do vídeo à medida que este é reproduzido e direcionar a atenção a um ponto principal? Analisar a sequência de exibição dos vídeos não se mostra válida neste caso, tendo em vista a aleatoriedade das listas. Por outro lado, na tentativa de responder a essa pergunta, também se agruparam os valores por período de

interesse em cada condição experimental, já que a sequência desses é comum em ambas as condições.

Na condição 1, observa-se um aumento da duração média de fixações na AI 1, à medida que se avança nos períodos de interesse, iniciando em 219,33 ms no período de interesse 1 e 262,06 ms no período de interesse 3. Ainda na mesma condição, também há aumento da média de fixações na AI 1 no período de interesse 2 (230,16 ms) em relação ao período de interesse 1 (229,42 ms), mas há uma redução considerável no período de interesse 3 (221,3 ms).

Na condição 2, o valor da duração média das fixações na AI 2 se manteve mais homogêneo entre os períodos de interesse, com números entre 237,08 ms e 247,84 ms. Na AI 1, por outro lado, há variação. A média para o período de interesse 1 foi de 169,75 ms, caindo para 88,91 ms no segundo período e 186,5 ms no terceiro período. Em relação à duração média de fixações, observa-se situação semelhante em relação a ambas as condições.

Na condição 1, na AI 1, o número médio de fixações aumenta, como na duração média, ao sair do período de interesse 1 para o período de interesse 2, mas reduz, ao se avançar para o período de interesse 3. Variando de 4,37 fixações no primeiro período, 5,38 no segundo e 5,44 no terceiro. Na AI 2, os valores seguem a mesma lógica, no entanto, na direção oposta, ou seja, reduzem do período 1 ao período dois e aumentam no período 3, com número médio de fixações de 2,25 no primeiro período de interesse, 2 no segundo e 2,6 no terceiro.

Na condição 2, na AI 2, o número médio de fixações também dialoga com a duração média de fixações nessa condição e área de interesse, ou seja, sendo inversamente proporcionais, com valores médios de 1 fixação no primeiro período de interesse, 1,5 fixações no segundo e 1,22 no terceiro, indicando que pode haver existência de fixações por parte de alguns sujeitos na AI 1. De forma semelhante, o número de fixações é maior na AI 2 e, assim como a duração média, também se mostram com características semelhantes.

Levando em consideração que a duração e número de fixações pode ser influenciado pelo tamanho das palavras ou velocidade das legendas, resume-se aqui as métricas previamente apresentadas na seção metodológica, considerando-se os períodos de interesse válidos para a análise. Ao se observar a densidade de informação de cada período, no tange ao tempo disponível para leitura, número de caracteres e número de caracteres por segundo, os valores

mostram-se semelhantes entre os três períodos de interesse, como resume a Tab.45, abaixo.

**Tabela 45** – Valores médio de duração, número de caracteres e caracteres/segundo por período de interesse

| Período de Interesse | Duração<br>(seg:frames) | Nº de caracteres | CPS   |
|----------------------|-------------------------|------------------|-------|
| P1                   | 01:33                   | 26               | 14,19 |
| P2                   | 01:27                   | 26,33            | 14,81 |
| P3                   | 01:24                   | 24,83            | 14,78 |

Fonte: elaborado pelo autor.

Tendo em vista que o fator de pesquisa se dá em torno da presença do texto em tela na AI 2, opta-se também por comparar os valores de número médio de fixações e da duração de fixações na condição 2. Esses dados ressoam de forma mais coerente com as métricas de tamanho médio das palavras (em caracteres) e a velocidade do texto em tela em (caracteres por segundo), apresentando diferenças na mesma direção.

Sendo assim, considerando os valores médios de fixações resumidos por Rayner (1998), a partir dos valores em ambas as condições, infere-se que, na condição 2, há melhor uniformidade em relação às durações médias de fixações na AI 2, onde há a presença de informação textual, bem como ao número médio de fixações. Ainda assim, há presença de fixações mais curtas na AI 1, sugerindo que pode ter havido busca de informações na região, mas que as fixações foram direcionadas à AI 2, de forma geral.

Na condição 1, o maior número médio de fixações na AI 1 pode sugerir que os participantes passaram a adotar a AI 1 como fonte de informação principal ou mais relevante em comparação ao número de fixações na AI 2, que, embora seja menor, apresenta ocorrência, sugerindo leitura no local. Tal observação também vai na linha de encontro ao se adicionar o número médio de fixações na AI, com valores médios crescentes do período 1 ao período 3. No entanto, as médias maiores referentes à AI 2 nos períodos 1 e 2, nessa condição, também demonstram atividade de leitura, porém, com número médio de fixações menores.

Ao se alinhar esses dois dados, infere-se que, inicialmente, o texto em tela foi considerado como fonte principal de informação antes de haver migração para a região de legenda, ou que a presença do texto em tela dispersa mais a atenção dos participantes, justificando as médias maiores com baixos números de fixações nessa região.

### **5.3 Tempo da Primeira Fixação na Área de Interesse**

Retomando a hipótese de pesquisa para essa medida, de modo geral, esperava-se que a primeira fixação fosse mais curta na AI 1 em comparação à AI 2, em detrimento da atenção dividida para a AI 2 na condição 1. De modo semelhante, por haver ausência de texto na AI 1, na condição 2, esperava-se que a primeira fixação na AI 2 fosse maior nessa condição, tendo em vista que seria o primeiro local de pouso do olhar quando o texto em tela aparecesse no vídeo. Contudo, ao se comparar apenas a AI 2 em ambas as condições, esperava-se que a média na condição 1 fosse maior do que na condição 2.

No agrupamento dos valores por participante, considerando ambas as condições, metade dos participantes (LPAP6, LDP16 e LDP20) apresentaram médias da primeira fixação mais longas na AI 2, com valores entre 210,09 ms para LPAP6 e 251, 72 ms para LDP20. E a outra metade (LDP11, LDP17 e LDP23) teve médias de primeira fixação mais longas na AI 1, com valores entre 154,18 ms para LDP23 e 215,09 ms para LDP11.

Sendo assim, no agrupamento geral por área de interesse, ao se observar o tempo médio da primeira fixação em cada AI, este foi maior na área de interesse AI 2, de forma geral, indo ao encontro da hipótese inicial até esse estágio de análise. Em parte, sugere-se que os valores superiores na AI 2 se dão devido, principalmente, às médias obtidas na condição total, já que a mesma situação é observada ao se analisar cada AI nas condições experimentais separadas, ou seja, os mesmo sujeitos apresentaram o comportamento ocular mencionado acima.

Os sujeitos LPAP6, LDP16 e LDP20 apresentaram média de primeira fixação mais longas na AI 2, com valores entre 174,5 ms e 288,71 ms. Os demais (LDP11, LDP17 e LDP23) apresentaram valores opostos, com primeiras fixações mais longas na AI 1, entre 146,55 ms e 209,55 ms, o que pode sugerir diferenças

individuais no processamento visual ou na estratégia de leitura nessa condição experimental.

Já na condição experimental 2, observa-se que as médias dos valores são relativamente mais homogêneas entre os participantes em comparação à condição 1, ao se comparar as duas áreas de interesse, ou seja, com médias de primeira fixação mais longas na AI 2 em quase todos os casos. Com exceção do participante LDP16, cuja média da primeira fixação se mostrou mais longa na AI 1 (186 ms) do que na AI 2 (149,22 ms), todos os demais participantes apresentaram médias mais longas na região do texto em tela.

A partir desse ponto, duas observações iniciais podem ser feitas. Primeiro, conforme era esperado, a duração da primeira fixação na AI 2 é levemente mais longa, assim como na condição 1, tendo em vista a presença da nova informação visual que aparece no vídeo na região central da tela. Segundo, os tempos de resposta da primeira fixação em AI 1 são mais curtos, já que na condição “legendas parciais”, a legenda não está presente na região inferior do vídeo, ficando visível apenas o texto em tela. Dessa forma, hipotetiza-se que, uma vez que o participante já venha com seu ritmo de leitura nessa região do vídeo, seria natural haver um atraso inicial à espera da próxima legenda, gerando os valores de fixações menores nesta área de interesse.

No agrupamento por período de interesse, os dados mostram uma mistura do que se observou na condição geral. Tanto na condição 1 quanto na condição 2, as médias mostraram durações mais longas na AI 2 em todos os três períodos de interesse. Os valores por período de interesse dialogam com as médias gerais por condição. Ao se comparar por área de interesse, assim como esperado na hipótese para essa variável, descrita na seção 3.8.2, a AI 2, na condição 1, apresentou média maior. Ao se comparar por condição, tanto a AI 1 quanto a AI 2 apresentaram médias de durações mais longas na condição 1, indicando uma possível maior demora de processamento nesta condição, como pontua Holmqvist *et al* (2011).

Levando em consideração a média de fixações gerais nas áreas de interesse, em ambas as condições, o tempo médio da primeira fixação mostrou-se menor do que a duração média de fixações em ambas as áreas de interesse, embora com menor diferença em relação à AI 2 na condição 1.

O teste estatístico de Mann-Whitney-Wilcoxon confirmou essa diferença

com um valor significativo ( $W = 1808$ ,  $p = 0.0075$ ), e o intervalo de confiança do rank (- 0.26; 95%CI [-0.43, -0.08]) reforça que a condição total demandou maior tempo para o processamento inicial da informação visual. Esse resultado sugere que a presença de legendas completas requer maior esforço cognitivo no momento inicial de fixação, provavelmente devido à maior quantidade de texto simultâneo competindo pela atenção do espectador. Por outro lado, a condição parcial, com menos informações textuais, parece facilitar o processamento inicial, resultando em tempos de fixação menores.

Contudo, vale ressaltar que, apesar de a média indicar tempo da primeira fixação levemente mais longo na AI 2, os valores encontram-se dentro das faixas consideradas esperadas para o processo de leitura, como resume Rayner (1998).

De forma geral, os valores das primeiras fixações na condição em ambas as condições se mostraram variáveis com indícios de comportamento de leitura mais individualizado, bem como a diferença nos valores ao se comparar as duas condições de pesquisa. Sendo assim, a seguir, também se apresenta a análise da sequência de fixações. Com essa medida, espera-se obter dados mais consistentes em relação ao padrão de leitura dos participantes no que tange à origem das fixações, e se elas ocorrem proeminentemente das mesmas áreas de interesse ou não.

#### **5.4 Análise de Sequência das Fixações**

Retomando as hipóteses para essa medida, esperava-se que, na condição 1, houvesse mais transições entre as áreas de interesse AI 1 e AI 2, com menor número de fixações originadas dentro na própria AI, e que, na condição 2, houvesse menos transições entre as áreas de interesse AI 1 e AI 2, com maior número de fixações originados da própria AI.

No agrupamento por participante, observou-se que, na condição 1, na AI 1, os participantes apresentam mais sequências de fixações originadas dentro da própria área de interesse, com valores entre 2,37 fixações para LDP16 e 4,33 para LDP17. Na AI 2, notou-se também a presença de algumas fixações vindas da área de interesse AI 1, embora em números menores e não ocorrendo na totalidade dos casos, como LDP17, cuja média de fixações na AI foi feita com sacadas dentro da própria AI.

Uma observação válida nesse ponto é que, como não se definiu uma área de interesse geral para a área do vídeo, não é possível, nessa etapa da análise, sugerir se essas fixações podem ter sido originadas de outros pontos de fixações. Além disso, a não ocorrência de fixações nessa análise não implica dizer que não tenha havido fixações na dada área de interesse. Pode ser um indício de baixas fixações ou fixações inexistentes, mas há a possibilidade de haver fixações originadas de outras regiões do vídeo, o que pode ser observado através de *scanpaths* (mapa com sequência de fixações e sacadas) ou gravação em vídeo das fixações de cada sujeito. Contudo, este não é o foco do trabalho, além de que investigações desse tipo seriam impraticáveis, considerando o número de condições de variáveis, totalizando 108 gravações para o dado experimento.

Na condição 2, a partir das médias individuais, observa-se que a maior sequência de fixações na AI 2 foi originada a partir da própria área de interesse, com valores entre 3,22 e 5,22 fixações, enquanto apenas 1 ou nenhuma fixação foi originada a partir da área de interesse AI 1. Em relação à AI 1, observa-se que as médias indicam que apenas uma ou nenhuma fixação foi originada da AI 2, como representam os valores “zero”, indicando que houve pouca transição entre as duas AIs durante a leitura das legendas.

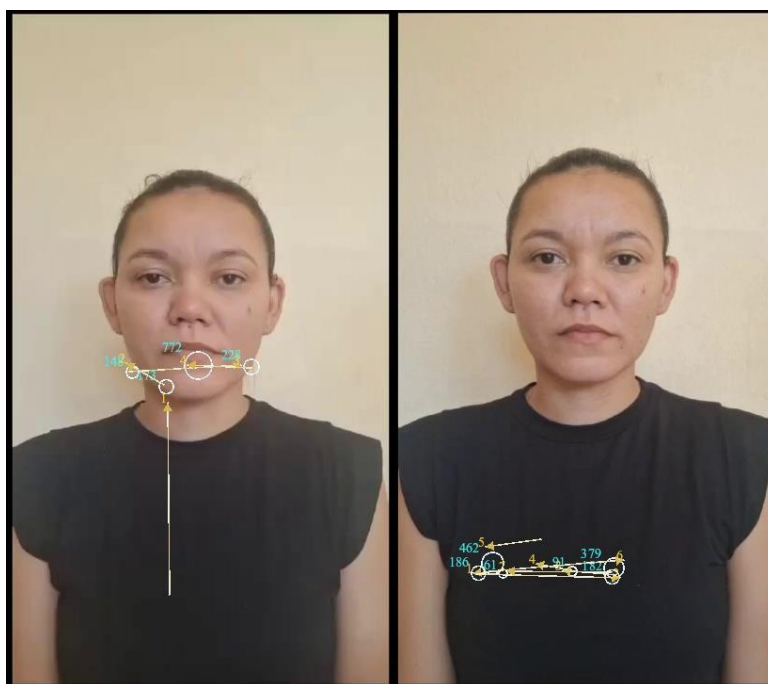
No cruzamento dos dados das duas condições, em relação às duas áreas de interesse vistas como predominantes em cada condição (AI para a condição 1 e AI 2 para a condição 2), embora os valores não apresentem diferenças consideráveis, nota-se um leve aumento de fixações médias na AI 2 na condição 2 (3,72 de 4,77), com fixações sequenciadas dentro da própria AI em comparação à AI 1 na condição 1 (3,56 de 4,75). Ou seja, não somente há mais fixações totais na condição 2, como também há mais fixações sequenciadas dentro da mesma AI, podendo indicar um processo de leitura mais fluído nessa condição.

Na sequência, são apresentadas as discussões referentes aos resultados do questionário pós-coleta.

## 5.5 Questionário Pós-Coleta

O questionário pós-coleta serviria de indício para o índice de compreensão do conteúdo apresentado nos vídeos, uma vez que as perguntas foram formuladas com base nas informações presentes nos textos em tela ou em segmentos de legendas anteriores ou posteriores a eles. Como apresentado na seção 4.6 dos seis participantes, apenas 2 apresentaram respostas equivocadas em relação às afirmativas propostas: LPD16 (afirmativas A5 e A11) e LDP20 (A10 e A11). Para entender melhor o padrão de visualização de ambos os participantes, recorreu-se à sequência de fixações visuais. A sequência de fixações de LPD16 em relação ao texto de recuperação de A5 pode ser visualizada na figura abaixo:

**Figura 38 – Sequência de Fixações de LPD16**



À esquerda, sequência de fixações referentes ao CP3, retiradas do vídeo 2 na condição 2. À direita, sequência referentes ao CP3, retiradas do vídeo 4 na condição 1.

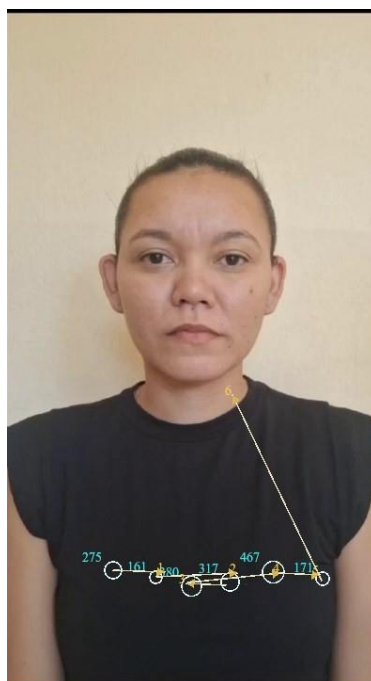
Fonte: elaborado pelo autor

No vídeo 2, observa-se que LPD16 sai da AI 1 após a leitura da legenda anterior, tem a primeira fixação computada na AI 2 na região mais central, com 173 ms de duração, direciona o olhar para o início do texto em tela, com uma fixação mais curta, vai para o final do texto e retorna para o centro da AI 2, com

uma fixação mais longa, de 772 ms. No vídeo 4, inicialmente, LDP16 tem a primeira fixação no início da legenda, direciona o olhar para a última palavra da legenda e retorna ao centro da AI 1, com uma saída de fixação da AI 1, e retorna ao início da legenda, finalizando a sequência de fixações mais um vez no centro da AI. Observa-se também fixações mais longas ao retornar para a AI 1, indicando mais tempo de demora de reengajamento na leitura dentro dessa AI. Além disso, nota-se que LDP16 não direciona o olhar para a AI 2, mantendo suas fixações na AI 1. A partir dos gráficos apresentados acima, LDP16 também teve tempo médio da primeira fixação maior em AI 2 em ambas as condições, além de média maior de duração na AI 2 na condição total.

Seguindo para LDP20, sua sequência de fixações para A11, no mesmo período de interesse, é apresentada na Fig. 39, a seguir.

**Figura 39** – Sequência de Fixações de LDP20 para CP3 no vídeo 4 na condição 1

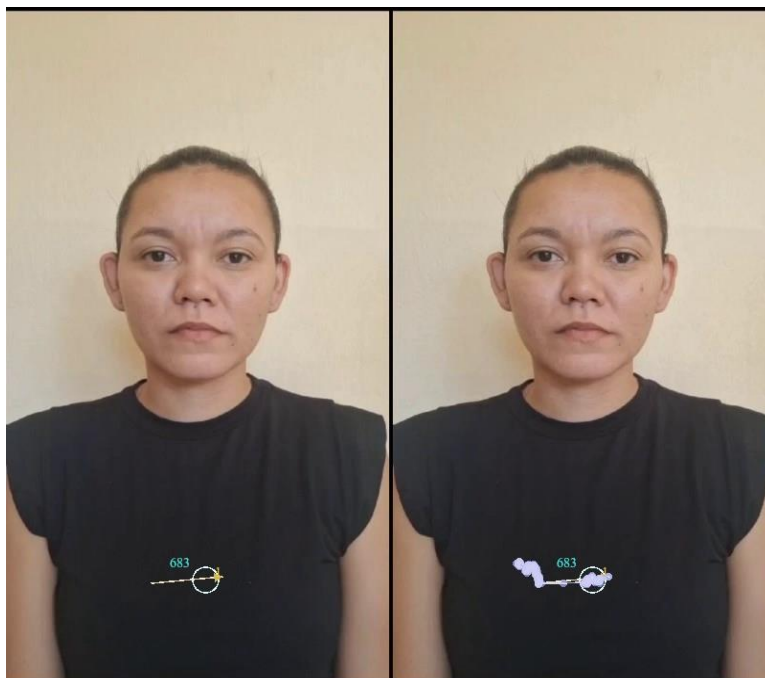


Fonte: elaborado pelo autor.

LDP20, em CP3, apresentou fixações mais homogêneas e direcionadas no percurso de fixações, com uma regressão após a segunda fixação. Na sequência, continua o percurso de fixações com sacadas visualmente homogêneas, finalizando seu trajeto ocular em direção à AI 2, cujo tempo de exibição já foi finalizado e, portanto, não permitindo nenhuma fixação nela. Em relação à AI 1,

sua sequência de fixações para LDP20 é apresentada na Fig. 40, a seguir.

**Figura 40** – Sequência de Fixações de LDP20 para CP1 no vídeo 4 na condição 1



À esquerda, mapa de fixação na faixa de 80 ms a 800 ms.  
À direita, amostra de posição do olhar antes da fixação computada.  
Fonte: elaborado pelo autor.

Diferentemente do movimento realizado em CP3 no mesmo vídeo, é possível ver, a partir do mapa de fixações por sequência, que não houve fixações válidas dentro da faixa de filtragem de dados suficientes para uma leitura completa do texto na área de interesse, o que pode sugerir a possível resposta incorreta para A10. Por outro lado, observa-se também que, apesar de o período de interesse estar na condição 2, também não houve fixações na AI 2.

De forma geral, os participantes acertaram as afirmativas de maneira satisfatória e, apesar das fixações levemente distribuídas entre legenda e texto em tela na condição 1, os dados indicam que não houve prejuízo na compreensão das afirmativas. Esses dados dialogam com os resultados encontrados por Josephson e Holmes (2006), cujas discussões pontuam que a adição dos textos em tela não prejudicou a compreensão do conteúdo.

Embora tenha havido 3 respostas incorretas na condição 1 em comparação à condição 2, com uma resposta incorreta, não se pode afirmar, a partir do ponto de vista dessa seção do trabalho, que a condição 1 possa ter influência negativa no entendimento, devido ao número insuficiente de dados.

Além disso, tanto LDP16 quanto LDP20 apresentaram respostas incorretas no mesmo vídeo, o que pode ter ocorrido em decorrência do conteúdo do vídeo, da forma como a afirmativa foi apresentada, especialmente na A11, ou outros fatores. Ainda assim, em face do número baixo de respostas incorretas e da baixa amostragem, não se pode obter conclusões somente com essa análise.

## 6 CONSIDERAÇÕES FINAIS

Este trabalho teve como objetivo geral investigar como o design de legendas em vídeos com texto em tela influencia o processamento de leitura. No capítulo 1, apresentou-se um panorama sobre a legendagem dentro dos Estudos da Tradução, mais especificamente no que tange à TAV. O capítulo 2 buscou delimitar o uso de legendas na presença de textos em tela, tanto em grandes produções disponíveis em plataformas de *streaming*, quanto na evolução desses textos inseridos em vídeos veiculados em redes sociais, como YouTube e Instagram. O capítulo exemplificou a ocorrência de textos em tela em produções audiovisuais, bem como o uso de legendas na presença de textos em tela em vídeos de redes sociais.

Dos exemplos, pôde-se observar que, dentro do mercado audiovisual tradicional, ou seja, materiais veiculados por grandes produtoras em plataformas de *streaming*, há diretrizes que pautam sobre o uso de legendas na presença de textos em tela tradicionais. No que tange aos vídeos nas redes sociais, observou-se a presença de dois padrões principais: o uso de legendas parciais e totais. No entanto, com a popularização das redes sociais, os exemplos analisados apontaram que não existem diretrizes para esses meios de veiculação no tange à produção de legendas, e cada produtor de conteúdo as utiliza de acordo com critérios próprios. O principal motivo se dá pela própria natureza dinâmica e interativa dessas plataformas, através das quais o público pode publicar seu próprio conteúdo, no intuito de não apenas divulgá-lo de forma geral, mas também promover produtos e marcas com fins de engajamento mercadológico. Sendo assim, o uso de textos dinâmicos, chamativos e de grande impacto se fazem cada vez mais presentes, assumindo um tipo de texto em tela.

Ainda no capítulo 2, apresentou-se um breve compilado de pesquisas sobre o uso de rastreamento ocular e legendagem em duas frentes: pesquisas convencionais com vídeos tradicionais legendados e pesquisas focadas em vídeos com predominância de textos na tela. Nessa seção, além das pesquisas de Liao, Kruger e Doherty (2020) e Liao et al. (2021), destacaram-se os trabalhos de Josephson e Holmes (2006) e O'Hagan e Sasamoto (2016), que fornecem uma descrição sobre o uso de texto em tela em noticiários, e Sasamoto (2024), que

também fornece discussão sobre textos em tela. A partir das discussões, observou-se que esses estudos investigam aspectos controláveis na produção de legendas que podem influenciar a recepção. No entanto, ficou evidente que há a necessidade de explorar melhor outros fatores em vídeos populares, como o uso do texto na tela de forma não regulamentada e como isso pode afetar a experiência do público. Dentre os fatores determinantes aqui, pontua-se a dificuldade de encontrar trabalhos com foco investigativo no uso de textos em tela com uso de legendas. Uma possível explicação pode se dar ao fato da não unificação de terminologia usada, tendo em vista que foram encontrados trabalhos que utilizam termos como “telop”, “legendas de impacto”, “legendas abusivas”.

No capítulo 3, no desenho experimental, privilegiou-se apresentar o processo escolhido para a construção da metodologia e das variáveis, a fim de analisar os vídeos legendados com a presença de textos em tela através do uso de rastreamento ocular. Embora os resultados sejam incipientes, ressalta-se aqui que a metodologia foi concebida no seguinte contexto: (1) escassez de pesquisas sobre rastreamento ocular em vídeos legendados com texto em tela; (2) trabalhos que trouxessem o controle das variáveis de análise; (3) a metodologia a ser usada com a movimentação ocular como ferramenta de análise. A falta de pesquisas com foco específico levou à elaboração de etapas de construção do experimento com vistas à experiência de outros trabalhos que também se utilizam do rastreamento ocular, a fim de obter variáveis controláveis, como número de palavras, velocidade das legendas, tamanho médio de palavras e frequência das palavras no léxico do português. Sem dúvida, a etapa de maior dificuldade está na construção do experimento no *Eye Link*, cujos modelos de pesquisa são predominantes para estudos com textos estáticos, com poucas orientações de fácil acesso para a construção de experimentos envolvendo materiais em vídeo.

Nos capítulos 4 e 5, foram apresentados e analisados os dados preliminares do estudo piloto com os seis participantes. A análise e discussão seguiu as medidas de investigação propostas na seção 2.8.2, tais quais: (1) o tempo de permanência na área de interesse, (2) médias de fixações por área de Interesse, tanto em termos de duração média quanto (3) do número de fixações, (4) o tempo da primeira fixação na área de interesse, (5) a análise da sequência de fixações e as respostas ao (6) questionário pós-coleta.

Retomando as hipóteses estipuladas para cada medida, para o tempo de permanência, esperava-se que, na condição 1, o tempo médio de permanência se mostrasse relativamente distribuído entre as duas AIs, e que, na condição 2, o tempo fosse maior na AI 2. Os resultados mostraram que, conforme esperado, o tempo médio de permanência foi maior na AI 1 na condição 1 e na AI 2 na condição 2, corroborando a hipótese de que a presença de texto influencia significativamente onde os espectadores focam sua atenção. Além disso, os dados podem indicar que, na condição 1, o texto em tela tende a dispersar a atenção dos participantes, justificando durações mais longas nessa área, com baixos números de fixações nessa região.

Semelhantemente, ao se comparar a duração média das fixações nas duas áreas de interesse, observou-se que a duração das fixações foi mais longa na área com texto ativo. Na condição 1, as fixações distribuíram-se equitativamente entre as legendas e os textos na tela, sugerindo uma divisão da atenção, embora com médias dentro do esperado para o padrão de leitura resumido por Rayner (1998). Já na condição 2, a fixação concentrou-se mais no texto da tela, dada a ausência de legendas na AI 1.

Para o tempo da primeira fixação na área de interesse, esperavam-se médias maiores na AI 2, tanto na condição 1 quanto na condição 2, mas, ao se comparar as duas AIs, esperava-se que a AI 2 na condição 2 tivesse tempos menores do que a duração média na AI 2 na condição 1. Os dados apresentados mostraram que os participantes demoraram mais tempo para fixar visualmente nas áreas de interesse na condição 1 em comparação com a condição 2, sugerindo uma interação dinâmica entre a presença de legendas e o texto na tela, ligada à rapidez com que os espectadores direcionaram a atenção ao visualizar conteúdo audiovisual. Sendo assim, os dados sugerem que o design das legendas na condição 1 pode ter exigido uma maior atenção.

Em relação à análise da sequência de fixações, para essa medida, esperava-se que, na condição 1, houvesse mais transições entre as áreas de interesse AI 1 e AI 2, com menor número de fixações originadas dentro na própria AI, e que, na condição 2, houvesse menos transições entre as áreas de interesse AI 1 e AI 2, com maior número de fixações originadas da própria AI. Os resultados indicaram que, na condição 1, os participantes mostraram mais transições entre as duas AIs, indicando uma distribuição mais equilibrada das

fixações, o que sugere uma tentativa de integrar as informações apresentadas simultaneamente em ambas as áreas. Já na condição 2, com a presença predominante de informação textual em uma única AI (AI 2), observou-se um número maior de fixações originadas e mantidas dentro dessa mesma área, refletindo uma concentração mais alta de atenção onde o conteúdo estava ativamente presente. Esse padrão sugere uma adaptação dos participantes à ausência de texto na AI 1, com foco maior na AI 2, onde havia texto em tela visível, destacando uma estratégia adaptativa ao design de visualização modificado e, possivelmente, uma leitura mais fluida.

Por fim, no que tange ao questionário pós-coleta, a maioria dos participantes demonstrou uma compreensão elevada do conteúdo dos vídeos, com quatro dos seis participantes acertando todas as afirmativas. Os dois participantes que não alcançaram pontuação máxima tiveram dificuldades específicas em afirmativas relacionadas a segmentos particulares dos vídeos. A análise das sequências de fixações desses participantes indicou que o desempenho pode ter sido influenciado pela maneira como as legendas e os textos em tela foram apresentados. Contudo, o baixo número dos dados não permite conclusões acerca do impacto de informações redundantes.

A partir dos resultados obtidos, retoma-se a problemática central que motivou este estudo: “Como o design de legendas em vídeos com texto em tela influencia a atenção visual e a experiência de leitura do espectador?” Ao longo desta pesquisa, foram delineadas hipóteses que orientaram tanto a estrutura metodológica quanto a análise dos dados. Presumia-se que o design de legendas parciais, ao evitar redundâncias entre textos em tela e legendas, promoveria uma distribuição mais uniforme da atenção entre as áreas de interesse e proporcionaria uma experiência de leitura visual mais fluida e integrada. Por outro lado, previa-se que o design de legendas totais, ao apresentar simultaneamente informações redundantes nas legendas e no texto em tela, aumentaria a competição por atenção, podendo afetar os padrões de leitura e, possivelmente, a compreensão.

Os resultados obtidos indicaram que, na condição 1, os participantes tenderam a concentrar sua atenção predominantemente na área das legendas, mesmo quando informações visuais equivalentes estavam disponíveis no texto em tela. Esse comportamento corrobora pesquisas anteriores que indicam que as legendas frequentemente funcionam como a principal fonte de informação em

vídeos. Contudo, ao contrário do previsto, essa redundância de informações não comprometeu a compreensão global dos vídeos. Ainda assim, notou-se que as transições entre áreas de interesse foram mais frequentes nesses casos, sugerindo maior demanda visual para processar informações redundantes apresentadas simultaneamente.

Na condição 2, o design de legendas parciais aparentou ser mais eficiente em promover uma experiência de leitura integrada. Com a remoção de informações redundantes nas legendas, os participantes apresentaram um comportamento visual mais estável, com menor número de transições entre as áreas de interesse e maior fluidez na leitura. Esse resultado sugere que as legendas parciais podem contribuir para otimizar a interação visual entre diferentes elementos verbais em tela, reduzindo potenciais sobrecargas visuais e promovendo maior conforto durante a leitura.

Em resumo, acredita-se que o trabalho cumpre com o que se propõe, embora apresente alguns pontos de melhoria a serem considerados futuramente do ponto de vista teórico-metodológico. Primeiro, como previamente apontado, a dificuldade em encontrar teoria existente sobre o uso de texto em tela e uso de legenda em rastreamento ocular, levando a uma base teórico limitada, na medida em que foi necessário buscar apoio em outras vertentes, aliando-se pesquisa em legendagem no geral com pesquisas isoladas sobre texto em tela, dificultando a conexão com teorias já existentes. Tal ponto levou a uma análise predominantemente em medidas descritivas. Além disso, o baixo número de participantes não permitiu alinhar análise quantitativa estatística para explorar as relações entre as variáveis. Por outro lado, espera-se que as discussões suscitadas aqui possam servir de base para a elaboração de teorias com foco específico no tempo que foi proposto.

Segundo, essas limitações de base teórica levaram a dificuldades para a elaboração do desenho experimental, que precisou adaptar metodologias não diretamente relacionadas, a fim de obter medidas aplicáveis no presente estudo. Por outro lado, ressalta-se a contribuição trazida em relação ao controle das variáveis, como velocidade de legenda, frequência e tamanho das palavras, cor, tamanho de fonte e posição das legendas nos vídeos produzidos. Como contribuição futura, espera-se adicionar outras variáveis ao estudo, a fim de se investigar a influência dessas, como velocidade de legenda, número de linhas e

design dos textos em tela e legendas nos vídeos.

Por fim, o *feedback* em relação aos participantes. A metodologia não se propõe a coletar *feedback* qualitativo dos participantes sobre a experiência de visualização. Tal *feedback* poderia oferecer *insights* adicionais sobre como os participantes perceberam as legendas e o conteúdo do vídeo, complementando os dados obtidos a partir do rastreamento ocular. Contudo, ressalta-se a descrição de uma metodologia detalhada com sequências que possam ser replicadas/adaptadas em pesquisas e construções metodológicas futuras, contribuindo, assim, para o desenvolvimento de técnicas e controle de variáveis nesse tipo de pesquisa experimental em legendagem.

Em resumo, embora os dados sejam considerados insuficientes do ponto de vista estatístico, há uma inclinação dos resultados para que as legendas na condição 2 possam favorecer um processo de leitura mais confortável para os espectadores. Sendo assim, os dados obtidos aqui podem destacar a importância de considerar como as legendas são apresentadas em materiais audiovisuais, especialmente em contextos com informações redundantes, a fim de que as informações sejam apresentadas de forma confortável.

## REFERÊNCIAS

1 MINUTE TALK. **What Girls Mean When They Te**, 2022, (59seg.). Disponível em: <https://www.youtube.com/shorts/gnWCOU0JUnw>. Acesso em: 22 ago. 2023.

3 ERRORES que TODO motociclista comete (5:49 min.). **FortNine en Español**, 2023. Disponível em: <https://youtu.be/ypw0n-LlI9w?t=204>. Acesso em 23 de jun. 2024.

AIRTABLE. [Internet]. San Francisco, CA: Airtable, 2023. Disponível em: <https://www.airtable.com/>. Acesso em: 19 out. 2023.

ANMARKRUD, Øistein; ANDRESEN, Anette; BRÅTEN, Ivar. Cognitive load and working memory in multimedia learning: Conceptual and measurement issues. **Educational Psychologist**, [local] v. 54, n. 2, p. 61-83, 2019.

ANTHONY, Laurence. **AntConc**. Version 4.2.4. Tokyo: Waseda University, 2023. Computer Software. Disponível em: <https://www.laurenceanthony.net/software>. Acesso em: 02 nov. 2023.

ARAÚJO, Vera Lúcia Santiago. **O processo de legendagem no Brasil**. 2002. Disponível em: [https://repositorio.ufc.br/bitstream/riufc/47224/1/2002\\_art\\_vlsaraujo.pdf](https://repositorio.ufc.br/bitstream/riufc/47224/1/2002_art_vlsaraujo.pdf). Acesso em: 24 jul. 2022.

ARAÚJO, Vera Lúcia Santiago; ASSIS, Ítalo Alves Pinto de. A segmentação linguística na legendagem para surdos e ensurdecidos (LSE) de 'Amor Eterno Amor': uma análise baseada em corpus. **Letras e Letras**, Uberlândia, v. 30, n. 2, p. 156–184, 2014.

ASSIS, Ítalo Alves Pinto de. **A influência do número de linhas e da velocidade no processamento de legendas por surdos e ouvintes: um estudo experimental com rastreador ocular**. Tese (Doutorado em Linguística Aplicada) – Universidade Estadual do Ceará, Fortaleza, 2021. Disponível em: [https://www.uece.br/posla/wp-content/uploads/sites/53/2021/07/TESE\\_%C3%8DTALO-ALVES-PINTO-DE-ASSIS.pdf](https://www.uece.br/posla/wp-content/uploads/sites/53/2021/07/TESE_%C3%8DTALO-ALVES-PINTO-DE-ASSIS.pdf). Acesso em: 15 jul. 2023.

BRASIL, Cartoon Network. ANIME BRASIL | FÃ CLUBE | CARTOON NETWORK (36 min). **Cartoon Network Brasil**, 2023. Disponível em: <https://youtu.be/FI-uqRcJ-0g?t=704>. Acesso em: 23 de jun. 2024.

BRASILEÑO prueba dulces mexicanos por primera vez! (20:06 min.). **Felipe Neto en Español**, 2024. Disponível em: <https://youtu.be/mMkKfXCBgzY?t=37>. Acesso em: 23 jun. 2024.

**BRAZILIAN Portuguese Timed Text Style Guide**. Parnet Help Center. NETFLIX, INC, 2024. Disponível em: <https://partnerhelp.netflixstudios.com/hc/enus/articles/215600497-Portuguese-Brazil-Timed-Text-Style-Guide>. Acesso em: 24 jun 2024.

CAFFREY, Colm. **Relevant abuse?:** Investigating the effects of an abusive subtitling procedure on the perception of TV anime using eye tracker and questionnaire. 2009. Tese (Doutorado) – Dublin City University, Dublin, 2009. Disponível em: <https://doras.dcu.ie/14835/>. Acesso em: 18 jun. 2024.

CARVALHO, Sandro Rogério Silva de; SEOANE, Alexandra Frazão. A relação entre unidades de velocidade de leitura na Legendagem para Surdos e Ensurdecidos. **Transversal - Revista em Tradução**, Fortaleza (CE), v. 5, n. 9, p. 137-153, 2019. Disponível em: [https://repositorio.ufc.br/bitstream/riufc/46883/1/2019\\_art\\_srscarvalhoafseoane.pdf](https://repositorio.ufc.br/bitstream/riufc/46883/1/2019_art_srscarvalhoafseoane.pdf). Acesso em: 18 Jun. 2024.

COJEAN, Salomé; MARTIN, Nicolas. Reducing the split-attention effect of subtitles during video learning: might the use of occasional keywords be an effective solution?. **L'Année Psychologique**, Paris, v. 121, n. 4, p. 417-442, 2021.

CONHEÇA a história da Regiane, que decidiu colocar o Google ao seu lado para se formar em TI. **Google Brasil**, 2023. Disponível em: [https://www.instagram.com/p/CzmGoTnALhS/?utm\\_source=ig\\_web\\_copy\\_link](https://www.instagram.com/p/CzmGoTnALhS/?utm_source=ig_web_copy_link). Acesso em: 24 nov. 2023.

DIAZ CINTAS, Jorge; ANDERMAN, Gunilla (ed.). **Audiovisual translation: Language transfer on screen**. Springer, 2008.

DÍAZ CINTAS, Jorge; REMAEL, Aline. **Audiovisual Translation: Subtitling. Translation Practices Explained**. Routledge: 2007.

DÍAZ CINTAS, Jorge; REMAEL, Aline. **Subtitling: Concepts and Practices. Translation Practices Explained**. Routledge: 2021. Disponível em: <https://doi.org/10.1093/biomet/52.3-4.591>. Acesso em: 10 set. 2024.

CHAVES, E. G. **Legendagem para surdos e ensurdecidos: um estudo baseado em corpus da segmentação nas legendas de filmes brasileiros em DVD**. 2012. Dissertação (Mestrado em Linguística Aplicada) - Universidade Estadual do Ceará, Fortaleza, 2012.

D'YDEWALLE, Gery; DE BRUYCKER, Wim. Eye movements of children and adults while reading television subtitles. **European Psychologist**, Newburyport, v. 12, n. 3, p. 196–205, 2007. DOI: <https://doi.org/10.1027/1016-9040.12.3.196>. Acesso em

EYELINK Data Viewer (version 4.4.1) [Computer software]. Oakville, Ontario, Canada: SR Research Ltd., 2024.

ELE nunca come depois do meio-dia. **NasDaily Português**, 2024, [s.n]. Disponível em: <https://www.youtube.com/shorts/ckEB0vhFpwl>. Acesso em: 26 maio. 2024.

EXPERIMENTANDO Doces Mexicanos! (28:28 min). **Felipe**

**Neto**, 2023. Disponível em: <https://youtu.be/dNsCf5o6PaY?t=48>. Acesso em: 23 de jun. 2024.

ENGLISH Timed Text Style Guide. Netflix.. N.D. Disponível em: <https://partnerhelp.netflixstudios.com/hc/en-us/articles/217350977-English-Timed-Text-StyleGuide>. Acesso em: 22 jun. 2024.

FLORAX, Mareike; PLOETZNER, Rolf. What contributes to the split-attention effect? The role of text segmentation, picture labelling, and spatial proximity. **Learning and instruction**, v. 20, n. 3, p. 216-224, 2010.

FRANCO, Eliana Paes Cardoso; ARAUJO, Vera Lúcia Santiago. Questões terminológico-conceituais no campo da tradução audiovisual (TAV). **Tradução em Revista**, v. 2021, n. 30, 2012. Disponível em: <https://www.maxwell.vrac.pucrio.br/18884/18884.PDF>. Acesso em: 16 jul. 2022.

GAMBIER, Yves. Introduction: Screen transadaptation: Perception and reception. **The Translator**, v. 9, n. 2, p. 171-189, 2003. Disponível em: <http://dx.doi.org/10.1080/13556509.2003.10799152>. Acesso em: 22 jun. 2024.

GODFROID, Aline. **Eye tracking in second language acquisition and bilingualism: A research synthesis and methodological guide**. Routledge, 2019.

GONÇALVES, Deborah Teles de Meneses. **O uso das mídias audiovisuais gratuitas como recurso emancipador para a obtenção autônoma da proficiência em língua inglesa**. São Cristóvão, 2021. Monografia (Licenciada em Letras Português-Inglês) – Departamento de Letras Estrangeiras, Centro de Educação e Ciências Humanas, Universidade Federal de Sergipe, São Cristóvão, SE, 2021

HOLMQVIST, Kenneth et al. **Eye tracking: A comprehensive guide to methods and measures**. oup Oxford, 2011.

JENSEMA, Carl J., et al. Eye Movement Patterns of Captioned Television Viewers. **American Annals of the Deaf**, [s.l], vol. 145, no. 3, pp. 275–85. *JSTOR*. 2000. Disponível em: <http://www.jstor.org/stable/44393197>. Acesso em: 14 abril 2024.

JAKOBSON, Roman. Aspectos linguísticos da tradução. Trad. Izidoro Blikstein. In: JAKOBSON, R, R. **Linguística e Comunicação**. São Paulo: Cultrix, p.63-86, 1995.

JOHNSON, Daniel. Polyphonic/pseudo-synchronic: Animated writing in the comment feed of Nicovideo. **Japanese Studies**, v. 33, n. 3, p. 297-313, 2013. Disponível em: <https://www.tandfonline.com/doi/full/10.1080/10371397.2013.859982>. Acesso em 24 mar. 2025.

JOSEPHSON, Sheree; HOLMES, Michael E. Clutter or content? How on-screen enhancements affect how TV viewers scan and what they learn. In: **Proceedings of the 2006 symposium on Eye tracking research & applications**. 2006. p. 155-162. Disponível em: <https://dl.acm.org/doi/abs/10.1145/1117309.1117361>.

Acesso em: 19 jun. 2024.

KRUGER, Jan-Louis et al. Legendas na imagem fílmica: uma visão geral dos estudos em rastreamento ocular. **Cadernos de Tradução**, v. 40, n. 2, p. 377-409, 2020. Disponível em:

<https://periodicos.ufsc.br/index.php/traducao/article/view/21757968.2020v40n2p377/43277>. Acesso em 25 mar. 2025.

KRUGER, Jan-Louis; SZARKOWSKA, Agnieszka; KREJTZ, Izabela. Subtitles on the moving image: An overview of eye tracking studies. **Refractory: a journal of entertainment media**, v. 25, p. 1-14, 2015. Disponível em:

<https://researchers.mq.edu.au/en/publications/subtitles-on-the-moving-image-an-overview-of-eye-tracking-studies>. Acesso em 25 mar. 2025.

KRUGER, Jan-Louis; WISNIEWSKA, Natalia; LIAO, Sixin. Why subtitle speed matters: Evidence from word skipping and rereading. **Applied Psycholinguistics**, [s.l.], v. 43, n. 1, p. 211-236, 2022.

<https://www.cambridge.org/core/journals/applied-psycholinguistics/article/why-subtitle-speed-matters-evidence-from-word-skipping-and-rereading/A7DB0CBF5910353143738CF9C1721317>. Acesso em 25 mar. 2025.

KRUGER, Jan-Louis; HEFER, Esté; MATTHEW, Gordon. Attention distribution and cognitive load in a subtitled academic lecture: L1 vs. L2. **Journal of Eye Movement Research**, [s.l.], v. 7, n. 5, 2014. Disponível em:

<https://researchmanagement.mq.edu.au/ws/portalfiles/portal/62321908/Publisher+version.pdf>. Acesso em: 18 jun. 2024.

KRUGER, Jan-Louis; STEYN, Faans. Subtitles and eye tracking: Reading and performance. **Reading Research Quarterly**, v. 49, n. 1, p. 105-120, 2014.

Disponível em: <https://www.jstor.org/stable/pdf/43497639.pdf>. Acesso em 25 mar. 2025.

KRUGER, Jan-Louis. Subtitles in the classroom: Balancing the benefits of dual coding with the cost of increased cognitive load. **Journal for Language Teaching= Ijenali Yekufundzisa Lulwimi= Tydskrif vir Taalonderrig**, [s.l.], v. 47, n. 1, p. 29-53, 2013. Disponível em:

<https://www.ajol.info/index.php/jlt/article/view/95475>. Acesso em 25 mar. 2025.

KÜNZLI, Alexander; EHRENSBERGER-DOW, Maureen. Innovative subtitling: A reception study. **Methods and strategies of process research: Integrative approaches in translation studies**, p. 187-200, 2011. Disponível em:

<https://digitalcollection.zhaw.ch/handle/11475/2996>. Acesso em: 19 jun. 2024.

LÂNG, Juha. Subtitles vs. narration: The acquisition of information from visual-verbal and audio-verbal channels when watching a television documentary.

**Eyetracking and applied linguistics**, v. 2, p. 59-81, 2016. Disponível em:

<https://mapaccess.uab.cat/publications/book-chapter/subtitles-vs-narration-acquisition-information-visual-verbal-and-audio>. Acesso em 24 mar. 2025.

LEXPORBR. **Léxico do Português Brasileiro**, 2019. Disponível em:

<https://www.lexicodoportugues.com/>. Acesso em: 12 de Out. 2023.

LIAO, Sixin et al. Using eye movements to study the reading of subtitles in video. **Scientific Studies of Reading**, v. 25, n. 5, p. 417-435, 2021.

LIAO, Sixin et al. The impact of audio on the reading of intralingual versus interlingual subtitles: Evidence from eye movements. **Applied psycholinguistics**, v. 43, n. 1, p. 237-269, 2022. Disponível em: <https://www.cambridge.org/core/journals/applied-psycholinguistics/article/impact-of-audio-on-the-reading-of-intralingual-versus-interlingual-subtitles-evidence-from-eye-movements/25E571F737EBCB68FF992910E54949E7>. Acesso em 24 mar. 2025.

LIAO, Sixin; KRUGER, Jan-Louis; DOHERTY, Stephen. The impact of monolingual and bilingual subtitles on visual attention, cognitive load, and comprehension. **The Journal of Specialised Translation**, v. 33, p. 70-98, 2020. Disponível em: <https://mapaccess.uab.cat/publications/journal-article/impact-monolingual-and-bilingual-subtitles-visual-attention-cognitive>. Acesso em 24 mar. 2025.

LINS, Thais Mazotti; CAMPOS, Giovana Cordeiro. Reflexões sobre o processo de adaptações culturais no filme *Divertida mente*. **Trama**, v. 18, n. 45, p. 25-37, 2022. MINDVALLEY BRASIL. N/A. Mindvalley, 2020. Disponível em: <https://www.instagram.com/p/CFHu0yWitvd/>. Acesso. 25 jun. 2024.

MATTHEW, Gordon. The effect of adding same-language subtitles to recorded lectures for non-native, English speakers in e-learning environments. **Research in Learning Technology**, v. 28, 2020. Disponível em: <https://journal.alt.ac.uk/index.php/rlt/article/view/2340>. Acesso em 24 mar. 2025.

MUNDAY, Jeremy. **Introducing Translation Studies: Theories and Applications**. Oxon: Routledge, 2001.

NASCIMENTO, A. K. P.. **Convencionalidade nas legendas de efeitos sonoros na legendagem para surdos e ensurdecidos (LSE)**. 2018. 241 f. Tese (Doutorado em Estudos da Tradução). Programa de Pós-Graduação em Estudos da Tradução, Faculdade de Filosofia, Letras e Ciências Humanas Universidade de São Paulo. São Paulo, 2018.

NAVES, S. B. et al. **Guia para produções audiovisuais acessíveis**. Secretaria do Audiovisual, Ministério da Cultura, 2016. Disponível em: <https://www.gov.br/culturaviva/pt-br/biblioteca-cultura-viva/documentos-e-publicacoes/documentos/minc-guia-para-producoes-audiovisuais-acessiveis-com-audiodescricao-das-imagens-2016.pdf/view>. Acesso: 12 de ago. de 2024.

NETFLIX Non-Branded Delivery Specifications. **Netflix**, INC, version 9.3, 2023. Disponível em: <https://partnerhelp.netflixstudios.com/hc/en-us/articles/215148917-Full-Non-Branded-Delivery-Specification-v9-3>. Acesso em: 12 jan. 2024.

NUÑES, Ana Carolina Santos Nunes. **Tradução para Legendas: Evolução no**

Filme *Star Wars: Uma Nova Esperança*. Monografia. Universidade Católica de Santos, 2017.

O'HAGAN, Minako; SASAMOTO, Ryoko. Crazy Japanese subtitles? Shedding light on the impact of impact captions with a focus on research methodology. **Eyetracking and applied linguistics**, v. 2, p. 31, 2016. Disponível em: <https://mapaccess.uab.cat/publications/book-chapter/crazy-japanese-subtitles-shedding-light-impact-impact-captions-focus>. Acesso em: 12 jan. 2024.

O'HAGAN, Minako. Japanese TV entertainment: Framing humour with open caption telop. **Translation, humor and the media**, v. 2, p. 70-88, 2010. Disponível em: <https://journals.sagepub.com/doi/abs/10.1177/1527476416677099>. Acesso em: 12 jan. 2024.

ORERO, Pilar et al. Conducting experimental research in audiovisual translation (AVT): A position paper. **JosTrans: The Journal of Specialised Translation**, n. 30, p. 105-126, 2018. Disponível em: <https://discovery.ucl.ac.uk/id/eprint/10053347/>. Acesso em: 12 jan. 2024.

ORUBE GAME STUDIO. **Apresentamos a vocês o Mombo Combo Legacy**, Fortaleza, 2023. Instagram: @orubegamestudio. Disponível em: [https://www.instagram.com/reel/Cviblj0N4yX/?utm\\_source=ig\\_web\\_copy\\_link&igshid=MzRIODBiNWFIzA%3D%3D](https://www.instagram.com/reel/Cviblj0N4yX/?utm_source=ig_web_copy_link&igshid=MzRIODBiNWFIzA%3D%3D). Acesso em: 24 nov. 2023.

PARK, Joseph Sung-Yul. **Regimenting languages on Korean television**: Subtitles and institutional authority. 2009.

RAYNER, Keith. Eye movements in reading and information processing: Twenty years of research. **Psychological Bulletin**. v. 124, 372–422. 1998.

PARK, Joseph. Ideological Functions of Subtitles in Multimodal Media Texts: Impact Captioning on Korean Television. In: the4th **International Conference on Multimodality**: From Print to Interactive Digital Media: Technology, Multimodal Representation and Knowledge, Singapore Management University, Singapore, July. 2008.

*PRIME VIDEO*, Amazon. Brazilian Portuguese (Brazil) Localization Style and Technical Guide, 2024. Disponível em: [https://videocentral.amazon.com/home/help?topicId=GBKB422Q9GYC7DWE&ref\\_=avd\\_sup\\_GBKB422Q9GYC7DWE#GFFTS6X436NEZ6Y8](https://videocentral.amazon.com/home/help?topicId=GBKB422Q9GYC7DWE&ref_=avd_sup_GBKB422Q9GYC7DWE#GFFTS6X436NEZ6Y8). Acesso em: 25 jun. 2024.

RAYNER, Keith et al. The effects of frequency and predictability on eye fixations in reading: implications for the EZ Reader model. **Journal of Experimental Psychology: Human Perception and Performance**, v. 30, n. 4, p. 720, 2004. Disponível em: <https://pubmed.ncbi.nlm.nih.gov/15301620/>. Acesso em 27 maio 2024.

RAYNER, Keith; JUHASZ, Barbara J.; POLLATSEK, Alexander. Eye Movements During Reading. In: M. J. Snowling & C. Hulme (Eds.), **The science of reading**: A

*handbook*, p. 79–97, Springer, 2005.

RAYNER, Keith; SERENO, Sara C.; RANEY, Gary E. Eye movement control in reading: a comparison of two types of models. **Journal of Experimental Psychology: Human Perception and Performance**, v. 22, n. 5, p. 1188, 1996. Disponível em: <https://pubmed.ncbi.nlm.nih.gov/8865619/>. Acesso em: 27 maio 2024.

REID, H. Literature on the screen: subtitle translation for public broadcasting. In: BART, W.; D'HAEN, T. (Eds.). **Something understood**. Studies in Anglo-Dutch literary translation. Amsterdam: Rodopi, p. 97-107, 1990.

SCOTT, Kate. **Pragmatics online**. London: Routledge, 2022.

SASAMOTO, Ryoko. **Relevance and Text-on-Screen in Audiovisual Translation: The Pragmatics of Creative Subtitling**. Taylor & Francis, 2024.

SASAMOTO, Ryoko; DOHERTY, Stephen; O'HAGAN, Minako. The 'hookability' of multimodal impact captions: A mixed-methods exploratory study of Japanese TV viewers. **Translation, Cognition & Behavior**, v. 4, n. 2, p. 253-280, 2021.

SHAPIRO, S. S.; WILK, M. B. An analysis of variance test for normality (complete samples). **Biometrika**, Reino Unido, v. 52, n. 3-4, p. 591-611, 1965. Disponível em: <https://www.jstor.org/stable/2333709>. Acesso em: 27 maio 2024.

SASAMOTO, Ryoko; O'HAGAN, Minako; DOHERTY, Stephen. Telop, affect, and media design: A multimodal analysis of Japanese TV programs. **Television & New Media**, [s.l.], v. 18, n. 5, p. 427-440, 2017. Disponível em: <https://journals.sagepub.com/doi/10.1177/1527476416677099>. Acesso em: 27 maio 2024.

VIEIRA, P. **A influência da segmentação e da velocidade na recepção de legendas para surdos e ensurdecidos**. 248f. 2016. Tese (Doutorado em Linguística Aplicada) – Programa de Pós-Graduação em Linguística Aplicada, Universidade Estadual do Ceará, Fortaleza-CE, 2016.

SR RESEARCH. **EyeLink 1000 Plus**: Multiple Eye Tracking Solutions. SR Research LTD, 2022.

SZARKOWSKA, Agnieszka; GERBER-MORÓN, Olivia. Viewers can keep up with fast subtitles: Evidence from eye movements. **PloS one**, v. 13, n. 6, p. e0199331, 2018. Disponível em: <https://journals.plos.org/plosone/article?id=10.1371/journal.pone.0199331>. Acesso em: 18 Jun. 2024.

TIRUMALA, Lakshmi N.; YOUNGBLOOD. Ed. Captioning Social Media Video. In: **Journal of Public Relations Education**, v. 7, n. 1, 169-187, 2021. Disponível em: [https://aejmc.us/jpre/wp-content/uploads/sites/25/2021/06/JPRE-71-5.22.21\\_editKO.pdf#page=169](https://aejmc.us/jpre/wp-content/uploads/sites/25/2021/06/JPRE-71-5.22.21_editKO.pdf#page=169). Acesso: 27 maio 2024.

THREE Mistakes Every Motorcyclist Makes (5:45 min.). In: **FortNine**, 2018.

Disponível em: <https://youtu.be/QWPypiml5c?t=204>. Acesso em: 23 de jun. 2024.

UNDERSTANDING Forced Narrative Subtitles. Parnet Help Center. **NETFLIX**, INC, 2017. Disponível em: <https://partnerhelp.netflixstudios.com/hc/en-us/articles/217558918-Understanding-Forced-Narrative-Subtitles>. Acesso em: 12 jan. 2024.

WILLIAMS, Gareth Ford. **Online Subtitling Editorial Guidelines**, V1.1. London: BBC, 2009. Disponível em: [https://www.louisholder.com/uploads/1/8/7/4/18742850/3.\\_bbc\\_online\\_sub\\_editorial\\_guidelines\\_vs1\\_1.pdf](https://www.louisholder.com/uploads/1/8/7/4/18742850/3._bbc_online_sub_editorial_guidelines_vs1_1.pdf)>. Acesso em: 22 jun. 2024.

## **APÊNDICE A – QUESTIONÁRIO PRÉ-COLETA**

### **TCLE 1**

Prezado (a) Senhor (a), você está sendo convidado a participar da pesquisa “Movimentação Ocular Durante a Leitura de Legendas Intralinguísticas em Vídeos com Texto em Tela”.

Esta pesquisa tem como objetivo investigar o movimento ocular durante a exibição de legendas com texto em tela, e o projeto está sendo desenvolvido pelo estudante de Mestrado Alberto Holanda Pimentel Neto, sob orientação da Profa. Dra. Maria Cristina Micelli.

Os procedimentos que serão realizados nesta pesquisa incluem diferentes etapas:

A etapa 1 consiste na apresentação e resposta por parte do participante a este questionário pré-coleta, mediante aceitação deste TCLE (Termo de Consentimento Livre e Esclarecido),

A etapa 2 consiste em assistir a vídeos legendados presencialmente no Laboratório de Aquisição, Desenvolvimento e Processamento das Linguagem, onde sua movimentação ocular será gravada com uma câmera de alta precisão durante a exibição dos vídeos.

Na sequência, na etapa 3 desta pesquisa, você responderá a um breve questionário Pós-Coleta, composto de 15 perguntas relacionadas aos vídeos apresentados, encerrando o ciclo da presente pesquisa.

As etapas 2 e 3, juntas, devem totalizar uma duração entre 30 min e 40 min, contando todo o processo de exibição, calibração, apresentação das instruções e resposta ao questionário Pós-Coleta.

Solicitamos a sua colaboração para as atividades previstas neste estudo, como também sua autorização para apresentar os resultados deste estudo em eventos científicos e publicar em revista científica nacional e/ou internacional.

Por ocasião da publicação dos resultados, o nome do participante será mantido em sigilo absoluto. A plataforma utilizada fará com que os dados do participante sejam mantidos anônimos e protegidos.

Informamos que essa pesquisa não tem nenhum benefício financeiro ou material. Os únicos riscos ao participar do experimento serão mínimos, entre eles: possível tédio e desconforto com o tempo de teste ou frustração caso não consiga realizar algum procedimento durante o estudo.

Essas questões tentarão ser evitadas ou minimizadas pelo pesquisador e pela professora que acompanharão cada etapa. Em caso de dano pessoal causado pelos procedimentos deste estudo, você tem direito de solicitar indenizações legalmente estabelecidas.

Esclarecemos que a participação no estudo é voluntária e, portanto, o(a) senhor(a) não é obrigado(a) a fornecer as informações e/ou colaborar com as atividades solicitadas pelo Pesquisador.

Caso decida não participar do estudo, ou resolva a qualquer momento desistir do mesmo, não sofrerá nenhum dano, nem haverá modificação na assistência que vem recebendo na Instituição. Em qualquer etapa do estudo, os pesquisadores estarão sempre à sua disposição para qualquer esclarecimento que considerem necessários.

Considerando que fui informado(a) dos objetivos e da relevância do estudo proposto, de como será minha participação, dos procedimentos e riscos decorrentes deste estudo, declaro o meu consentimento em participar da pesquisa, como também concordo que os dados obtidos na investigação sejam utilizados para fins científicos (divulgação em eventos e publicações). As informações obtidas serão analisadas em conjunto com as de outros voluntários, não sendo divulgada a identificação de nenhum participante.

Eu tenho o direito de ser mantido atualizado sobre os resultados parciais das pesquisas, quando em estudos abertos, ou de resultados que sejam do conhecimento dos pesquisadores. Estou ciente que receberei uma via desse documento.

Os investigadores responsáveis são o estudante Alberto Holanda Pimentel Neto e a Profa. Dra. Maria Cristina Micelli Fonseca, que podem ser encontradas nos e-mails:

- albertoholanda2014@gmail.com
- mcristina.fonseca@ufc.br

Se você tiver alguma consideração ou dúvida sobre a ética da pesquisa, entre em contato com o Comitê de Ética em Pesquisa da UFC/PROPESQ, localizado na Rua Coronel Nunes de Melo, 1000 - Rodolfo Teófilo. Horário de segunda a sexta-feira, das 8h ao meio-dia. Se está de acordo, preencha abaixo seu nome e a data de hoje:

Está de acordo em participar da pesquisa?

☐ Estou de acordo e desejo participar da pesquisa

☐ Não estou de acordo e não desejo participar da pesquisa

Nome do Participante: \_\_\_\_\_

E-mail: \_\_\_\_\_

Telefone: \_\_\_\_\_

Sexo Biológico: ☐ Masculino ☐ Feminino Idade \_\_\_\_\_

Data de Nascimento:     /     /  
Estado de Nascimento: \_\_\_\_\_  
Formação: \_\_\_\_\_

Que nota você daria para sua leitura em Língua Portuguesa (1-10)?

Caracterize a sua atitude para a escola quando criança (1 – 5)

1 – odiava ir à escola

2 – adorava ir à escola

Sentiu dificuldades para aprender a ler na escola primária (1-5)?

1 – não apresentou nenhuma dificuldade

5 – apresentou muitas dificuldades

Você tem alguma dificuldade com a visão? [   ] Sim [   ] Não

Se marcou “sim”, qual dificuldade/problema de visão?

[   ] Miopia

[   ] Astigmatismo

[   ] Hipermetropia

[   ] Catarata

[   ] Daltonismo

[   ] Estrabismo

[   ] Glaucoma

[   ] Presbiopia

[   ] Outras(s)

Grau de Miopia/Hipermetropia/Astigmatismo\_OD

\_\_\_\_\_

Grau de Miopia/Hipermetropia/Astigmatismo\_OE

\_\_\_\_\_

Você usa óculos ou lentes de contato para a dificuldade com a visão?

[   ] Sim, uso óculos

[   ] Sim, uso lentes

[   ] Não.

Sua lateralidade é:

[   ] Destra (escreve com a mão direita)

[   ] Canhota (escreve com a mão esquerda)

Você tem algum problema neurológico?

[   ] Sim

[   ] Não

Caso tenha algum problema neurológico, por favor, especifique

\_\_\_\_\_

Você tem algum problema psiquiátrico, como, por exemplo, depressão, ansiedade, dentre outros?

☐ Sim

☐ Não

Caso tenha algum problema psiquiátrico, por favor, especifique

---

Você faz uso de remédios controlados ou drogas/substâncias que afetam o sistema nervoso central?

☐ Sim

☐ Não

Caso faça uso, por favor, especifique qual remédio/substância utiliza \_\_\_\_\_

Você costuma ver filmes ou programas de TV legendados?

☐ Sim

☐ Não

Com que frequência você costuma ver programas legendados?

☐ Diariamente

☐ 3 vezes por semana

☐ 2 vezes por semana

☐ 1 vez por semana

☐ raramente

Você tem alguma dificuldade em ver filmes/programas legendados?

☐ Sim

☐ Não

Por favor, especifique qual dificuldade você tem ao ver filmes/programas legendados

---

Você costuma acessar redes sociais, como Instagram, Facebook e YouTube?

☐ Sim

☐ Não

Com que frequência você costuma acessar as redes sociais?

☐ Diariamente

☐ 3 vezes por semana

☐ 2 vezes por semana

☐ 1 vez por semana

☐ raramente

Quantas horas, em média, você consome conteúdo nas redes sociais?

- ☐ Menos de 1 hora por dia
- ☐ Entre 1 e 3 horas por dia
- ☐ Entre 3 e 5 horas por dia
- ☐ Mais de 5 horas por dia

Qual(is) plataforma(s) você mais utiliza?

- ☐ Instagram
- ☐ Facebook
- ☐ YouTube
- ☐ Tiktok
- ☐ Twitter
- ☐ Kiwi
- ☐ Snapchat

Você costuma navegar nas redes sociais com o áudio ativado ou desativado?

- ☐ Áudio ativado
- ☐ Áudio desativado



9. Afirmativa 9 - O TDAH não é uma condição tratável.

☐ Verdadeiro ☐ Falso

10. Afirmativa 10 - Treinar até a falha deve ser feito apenas por atletas avançados.

☐ Verdadeiro ☐ Falso

11. Afirmativa 11 - O aquecimento lubrifica as articulações.

☐ Verdadeiro ☐ Falso

12. Afirmativa 12 - Nutrição, sono e aquecimento foram as principais dicas apresentadas em um dos vídeos.

☐ Verdadeiro ☐ Falso

13. Afirmativa 13 - É necessário priorizar a saúde mental.

☐ Verdadeiro ☐ Falso

14. Afirmativa 14 - Praticar exercícios físicos não foi listado como dica para qualidade de vida.

☐ Verdadeiro ☐ Falso

15. Afirmativa 15 - É necessário hidratar-se com frequência e ter uma alimentação balanceada.

☐ Verdadeiro ☐ Falso

## **APÊNDICE C—ROTEIRO PARA O VÍDEO 1**

Uma boa noite de sono é essencial para a nossa saúde e bem-estar. Aqui estão algumas dicas rápidas para melhorar a qualidade do seu sono: Dica 1: Tenha uma Rotina de Sono Consistente: Estabeleça um horário regular para dormir e acordar, mesmo nos fins de semana. Isso ajuda a sincronizar o seu relógio biológico, tornando mais fácil adormecer e acordar na hora desejada. Dica 2: Durma em um Ambiente Confortável: Seu quarto deve ser um lugar calmo, escuro e com uma temperatura agradável. Evite dispositivos eletrônicos antes de dormir, pois a luz pode interferir no sono. Opte por atividades relaxantes, como ler um livro. Dica 3: Mantenha uma Alimentação balanceada: Evite refeições pesadas e picantes antes de dormir. Abstenha-se de cafeína e álcool perto da hora de dormir, pois são estimulantes que podem atrapalhar o sono. Lembre-se de que pequenas mudanças nos seus hábitos diários podem resultar em uma grande melhoria na qualidade do seu sono. Comece hoje a implementar essas dicas e desfrute de noites de sono mais revigorantes.

## **APÊNDICE D—ROTEIRO PARA O VÍDEO 2**

Trabalhar em casa, no chamado “home office”, requer disciplina para manter a produtividade. Por isso, separei três dicas essenciais para ter sucesso. Defina um horário de Trabalho: Estabeleça um horário consistente no qual você seja mais produtivo. Descubra se você rende melhor de manhã, à tarde ou à noite e mantenha uma rotina que funcione para você. Vista-se de forma adequada: Vestir-se como se estivesse no escritório pode te ajudar a entrar no modo trabalho. Abandone o pijama. Roupas profissionais contribuem para a mentalidade de trabalho, além de serem ideais para videoconferências. Tenha um Espaço de Trabalho Específico: Crie um ambiente dedicado ao trabalho, livre de distrações. Tenha uma mesa organizada, uma cadeira confortável e mantenha seu espaço limpo. Isso ajuda a manter o foco. Lembre-se dessas dicas para otimizar seu home office e ser mais produtivo. Com um horário de trabalho bem definido, vestimenta adequada e espaço de trabalho organizado, você pode fazer do home office uma experiência eficiente e gratificante.

### **APÊNDICE E—ROTEIRO PARA O VÍDEO 3**

O Transtorno do Déficit de Atenção e Hiperatividade, conhecido como TDAH, merece atenção. Três sinais-chave podem ajudar a identificá-lo. Primeiro: a desatenção constante é uma característica marcante. Pessoas com TDAH têm dificuldade em manter o foco, cometem erros por falta de atenção e frequentemente se esquecem de compromissos importantes. Em segundo lugar, a hiperatividade e inquietude são comuns. Manifestam-se como necessidade constante de movimentar-se, dificuldade em permanecer sentado e agitação. Pessoas com TDAH também tendem a falar muito e interromper os outros, afetando suas relações. Além disso, a impulsividade é notável. Pessoas com TDAH têm dificuldade em controlar impulsos e tomam decisões precipitadas sem considerar as consequências. Para o diagnóstico de TDAH, esses sintomas devem ser persistentes e causar dificuldades na vida cotidiana. O tratamento envolve terapia, estratégias de gerenciamento e, às vezes, medicamentos. O TDAH não é uma fraqueza. É uma condição tratável. Se você ou alguém que você conhece apresenta esses sinais, busque ajuda de um profissional de saúde. O tratamento certo faz a diferença na qualidade de vida das pessoas com TDAH. Mantenha-se atento aos sinais.

## **APÊNDICE F—ROTEIRO PARA O VÍDEO 4**

Praticar exercícios físicos de forma eficaz requer algumas dicas essenciais. Treine até a falha. Não importa se você é iniciante ou um atleta avançado, treinar até a falha é crucial. Dê o seu melhor até não conseguir mais realizar uma repetição adequada. Essa abordagem desafia seus limites e estimula o progresso. Tenha a mentalidade correta: Desenvolver a mentalidade certa é a chave para o sucesso. Comprometa-se a dar o seu máximo, independentemente do seu nível. Ter essa mentalidade envolve esforço constante e um desejo de melhoria contínua. Priorize nutrição, sono e aquecimento. Além do treinamento, não negligencie a importância da nutrição. Mantenha uma dieta equilibrada, evite maus hábitos alimentares e o excesso de álcool. Priorize o sono, pois o descanso adequado é vital para um treinamento eficaz. E lembre-se do aquecimento. Ele prepara os músculos e lubrifica as articulações, prevenindo lesões. Investir em sua saúde física requer dedicação em todas essas áreas: treinamento, mentalidade, nutrição, sono e aquecimento. Essas dicas essenciais garantem que você alcance o melhor em seu condicionamento físico.

## **APÊNDICE G—ROTEIRO PARA O VÍDEO 5**

Viver bem deveria ser nossa meta número um. Portanto, neste vídeo, compartilharemos três dicas para melhorar sua saúde e bem-estar. Dica número um: Tenha uma dieta equilibrada, começando pela alimentação. Mantenha uma dieta equilibrada com frutas, vegetais, grãos integrais, proteínas magras e laticínios com baixo teor de gordura. Evite alimentos processados ricos em açúcar e gorduras saturadas. Dica número dois: Pratique exercícios. Exercite-se regularmente, com pelo menos 150 minutos de atividade aeróbica moderada ou 75 minutos de atividade intensa por semana. Varie suas atividades para manter o interesse e obter benefícios para a saúde física e mental. Dica número três: Priorize sua saúde mental. Pratique técnicas de gerenciamento de estresse, mantenha conexões sociais significativas e não hesite em buscar ajuda profissional se necessário. Em resumo, manter uma dieta equilibrada, praticar exercícios regulares e ter atenção à saúde mental são os pilares para o bem-estar. Lembre-se: Sua saúde é uma prioridade. Implemente essas mudanças simples para ter uma vida mais saudável e feliz.

## **APÊNDICE H—ROTEIRO PARA O VÍDEO 6**

Manter uma dieta saudável é fundamental para nossa saúde e bem-estar. Hoje, vamos compartilhar três dicas valiosas para melhorar sua alimentação. Primeira dica: Tenha uma alimentação balanceada. Consuma frutas, vegetais, grãos inteiros, proteínas magras e laticínios com baixo teor de gordura. Uma paleta colorida de alimentos garante uma ingestão equilibrada de nutrientes. Segunda dica: Alimente-se de forma Consciente. Evite distrações ao comer e concentre-se na refeição. Mastigar devagar ajuda a controlar as porções e a reconhecer os sinais de fome e saciedade. Terceira dica: Hidrate-se com frequência. Beba água regularmente para ajudar na digestão, absorção de nutrientes e eliminação de resíduos. Oito copos por dia são um bom ponto de partida, mas ajuste conforme sua necessidade. Lembre-se de que a consistência é fundamental. Faça mudanças progressivas em direção a uma dieta mais saudável. Seguindo essas dicas, você estará no caminho certo para uma vida mais saudável e equilibrada. Consulte um profissional de saúde ou nutricionista para orientações personalizadas. Cuide de sua saúde e bem-estar, fazendo escolhas alimentares conscientes.

## APÊNDICE I – TERMO DE CONSENTIMENTO DE USO DE IMAGEM E VOZ

Eu, \_\_\_\_\_, brasileira,  
solteira, portadora do CPF nº \_\_\_\_\_, RG  
nº \_\_\_\_\_, residente à  
Rua \_\_\_\_\_,

\_\_\_\_\_, 700, Bloco D,  
Apto. 402, no Bairro Barroso, na cidade de Fortaleza - CE, autorizo expressamente, para fins de divulgação acadêmica e científica, a captação, o uso, a guarda e a exibição/execução da minha imagem e voz, em caráter definitivo e gratuito, em toda e qualquer participação minha na proposta descrita abaixo, para fins exclusivamente educacionais. Estou ciente de que filmagens, fotos e capturas de tela decorrentes da minha participação na proposta podem ser utilizadas a qualquer tempo pela parte autorizada.

Proposta: Pesquisa Acadêmica a nível de Mestrado intitulada “**Movimentação Ocular Durante a Leitura de Legendas Intralinguísticas em Vídeos com Texto em Tela**”.

Aluno Responsável: Alberto Holanda Pimentel Neto  
Profa. Orientadora: Dra. Maria Cristina Micelli Fonseca

---

Samara Myria de Oliveira Paiva

---

Maria Cristina Micelli Fonseca

---

Alberto Holanda Pimentel Neto

Fortaleza - CE, DATA \_\_\_\_\_

## APÊNDICE J–TERMO DE CONSENTIMENTO LIVRE E ESCLARECIDO

Você está sendo convidado por **Alberto Holanda Pimentel Neto** a participar da pesquisa intitulada “**Movimentação Ocular Durante a Leitura de Legendas Intralinguísticas em Vídeos com Texto em Tela**”. O presente projeto está sendo realizado a partir da pesquisa de Pós-Graduação a nível de Mestrado no Programa de Pós-Graduação em Estudos da Tradução na Universidade Federal do Ceará, sob supervisão da Professora Dra. Maria Cristina Micelli Fonseca. Leia atentamente as informações abaixo e faça qualquer pergunta que desejar, para que todos os procedimentos desta pesquisa sejam esclarecidos. Você não deve participar contra a sua vontade.

Esta etapa da pesquisa prevê exibição de vídeos legendados no computador. O rastreador ocular estará gravando o movimento dos seus olhos ao assistir aos vídeos.

Em qualquer momento você pode se recusar a continuar participando da pesquisa, retirando o seu consentimento, sem que isso lhe traga qualquer prejuízo. Garantimos que as informações obtidas através da sua participação não permitirão a identificação da sua pessoa, exceto aos responsáveis pela pesquisa, e que a divulgação das mencionadas informações só será feita entre os profissionais estudiosos do assunto.

Endereço d(os, as) responsável(is) pela pesquisa:

**Nome:** Alberto Holanda Pimentel Neto  
**Instituição:** Universidade Federal do Ceará  
**Endereço:** Av. da Universidade, 2762, Dpto. de Letras Vernáculas, Lab. 4: Laboratório de Aquisição, Desenvolvimento e Processamento da Linguagem.  
**Telefones para contato:** (83) 99641-4607

**ATENÇÃO:** Se você tiver alguma consideração ou dúvida, sobre a sua participação na pesquisa, entre em contato com o Comitê de Ética em Pesquisa da UFC/PROPEQ – Rua Coronel Nunes de Melo, 1000 - Rodolfo Teófilo, fone: 3366- 8344/46. (Horário: 08:00-12:00 horas de segunda a sexta-feira). O CEP/UFC/PROPEQ é a instância da Universidade Federal do Ceará responsável pela avaliação e acompanhamento dos aspectos éticos de todas as pesquisas envolvendo seres humanos.

O abaixo assinado \_\_\_\_\_, \_\_\_\_\_ anos, RG: \_\_\_\_\_, declara que é de livre e espontânea vontade que está como participante de uma pesquisa. Eu declaro que li cuidadosamente este Termo de Consentimento Livre e Esclarecido e que, após sua leitura, tive a oportunidade de fazer perguntas sobre o seu conteúdo, como também sobre a pesquisa, e recebi explicações que responderam por completo minhas dúvidas. E declaro, ainda, estar recebendo uma via assinada deste termo.

Fortaleza,        /        /

Participante: \_\_\_\_\_  
Data: \_\_\_\_\_  
Assinatura: \_\_\_\_\_

Nome do pesquisador: \_\_\_\_\_  
Data: \_\_\_\_\_  
Assinatura: \_\_\_\_\_

Nome do profissional que aplicou o TCLE: \_\_\_\_\_  
Data: \_\_\_\_\_  
Assinatura: \_\_\_\_\_

### APÊNDICE K– MENSAGENS DA LISTA A

| Mensagem ID | Trial | Mensagem - IP time | Mensagem - Text ID |
|-------------|-------|--------------------|--------------------|
| MSG         | -1    | -8213              | CP1Start           |
| MSG         | -1    | -10319             | CP1End             |
| MSG         | -1    | -30771             | CP2Start           |
| MSG         | -1    | -32134             | CP2End             |
| MSG         | -1    | -42028             | CP3Start           |
| MSG         | -1    | -43555             | CP3End             |
| MSG         | -1    | -65599             | CP4Start           |
| MSG         | -1    | -66661             | CP4End             |
| MSG         | -4    | -12000             | CP1Start           |
| MSG         | -4    | -13366             | CP1End             |
| MSG         | -4    | -27198             | CP2Start           |
| MSG         | -4    | -29131             | CP2End             |
| MSG         | -4    | -46230             | CP3Start           |
| MSG         | -4    | -48030             | CP3End             |
| MSG         | -4    | -72944             | CP4Start           |
| MSG         | -4    | -74494             | CP4End             |
| MSG         | -6    | -14888             | CP1Start           |
| MSG         | -6    | -16599             | CP1End             |
| MSG         | -6    | -30791             | CP2Start           |
| MSG         | -6    | -32948             | CP2End             |
| MSG         | -6    | -55462             | CP3Start           |
| MSG         | -6    | -56928             | CP3End             |
| MSG         | -6    | -62962             | CP4Start           |
| MSG         | -6    | -64562             | CP4End             |
| MSG         | -10   | -11199             | CP1Start           |
| MSG         | -10   | -13332             | CP1End             |
| MSG         | -10   | -30697             | CP2Start           |
| MSG         | -10   | -32264             | CP2End             |

|     |     |        |          |
|-----|-----|--------|----------|
| MSG | -10 | -49729 | CP3Start |
| MSG | -10 | -51929 | CP3End   |
| MSG | -10 | -71860 | CP4Start |
| MSG | -10 | -73726 | CP4End   |
| MSG | -12 | -7299  | CP1Start |
| MSG | -12 | -8832  | CP1End   |
| MSG | -12 | -24098 | CP2Start |
| MSG | -12 | -26164 | CP2End   |
| MSG | -12 | -61261 | CP3Start |
| MSG | -12 | -63228 | CP3End   |
| MSG | -12 | -79227 | CP4Start |
| MSG | -12 | -80827 | CP4End   |
| MSG | -16 | -14300 | CP1Start |
| MSG | -16 | -17149 | CP1End   |
| MSG | -16 | -30465 | CP2Start |
| MSG | -16 | -32565 | CP2End   |
| MSG | -16 | -44297 | CP3Start |
| MSG | -16 | -45930 | CP3End   |
| MSG | -16 | -77727 | CP4Start |
| MSG | -16 | -79861 | CP4End   |

### APÊNDICE L – MENSAGENS DA LISTA D

| Mensagem ID | Trial | Mensagem - IP time | Mensagem - Text ID |
|-------------|-------|--------------------|--------------------|
| MSG         | -1    | -11199             | CP1Start           |
| MSG         | -1    | -13332             | CP1End             |
| MSG         | -1    | -30697             | CP2Start           |
| MSG         | -1    | -32264             | CP2End             |
| MSG         | -1    | -49729             | CP3Start           |
| MSG         | -1    | -51929             | CP3End             |
| MSG         | -1    | -71860             | CP4Start           |
| MSG         | -1    | -73726             | CP4End             |
| MSG         | -4    | -14888             | CP1Start           |
| MSG         | -4    | -16599             | CP1End             |
| MSG         | -4    | -30791             | CP2Start           |
| MSG         | -4    | -32948             | CP2End             |
| MSG         | -4    | -55462             | CP3Start           |
| MSG         | -4    | -56928             | CP3End             |
| MSG         | -4    | -62962             | CP4Start           |
| MSG         | -4    | -64562             | CP4End             |
| MSG         | -10   | -7299              | CP1Start           |
| MSG         | -10   | -8832              | CP1End             |
| MSG         | -10   | -24098             | CP2Start           |
| MSG         | -10   | -26164             | CP2End             |
| MSG         | -10   | -61261             | CP3Start           |
| MSG         | -10   | -63228             | CP3End             |
| MSG         | -10   | -79227             | CP4Start           |
| MSG         | -10   | -80827             | CP4End             |
| MSG         | -12   | -12000             | CP1Start           |

|     |     |        |          |
|-----|-----|--------|----------|
| MSG | -12 | -13366 | CP1End   |
| MSG | -12 | -27198 | CP2Start |
| MSG | -12 | -29131 | CP2End   |
| MSG | -12 | -46230 | CP3Start |
| MSG | -12 | -48030 | CP3End   |
| MSG | -12 | -72944 | CP4Start |
| MSG | -12 | -74494 | CP4End   |
| MSG | -15 | -8213  | CP1Start |
| MSG | -15 | -10319 | CP1End   |
| MSG | -15 | -30771 | CP2Start |
| MSG | -15 | -32134 | CP2End   |
| MSG | -15 | -42028 | CP3Start |
| MSG | -15 | -43555 | CP3End   |
| MSG | -15 | -65599 | CP4Start |
| MSG | -15 | -66661 | CP4End   |
| MSG | -17 | -14888 | CP1Start |
| MSG | -17 | -16599 | CP1End   |
| MSG | -17 | -30791 | CP2Start |
| MSG | -17 | -32948 | CP2End   |
| MSG | -17 | -55462 | CP3Start |
| MSG | -17 | -56928 | CP3End   |
| MSG | -17 | -62962 | CP4Start |
| MSG | -17 | -64562 | CP4End   |

## ANEXO A–PARECER COMITÊ DE ÉTICA

UNIVERSIDADE FEDERAL DO  
CEARÁ PROPESQ - UFC



### PARECER CONSUBSTANCIADO DO CEP

#### DADOS DO PROJETO DE PESQUISA

**Título da Pesquisa:** Movimentação Ocular durante a leitura de legendas intralinguísticas em vídeos com texto em tela

**Pesquisador:** ALBERTO HOLANDA PIMENTEL NETO

**Área Temática:**

**Versão:** 2

**CAAE:** 80810724.0.0000.5054

**Instituição Proponente:** Universidade Federal do Ceará/ PROPESQ

**Patrocinador Principal:** Financiamento Próprio

#### DADOS DO PARECER

**Número do Parecer:** 7.045.257

#### Apresentação do Projeto:

A legendagem, no que tange os Estudos da Tradução, é um tipo de tradução audiovisual, uma vez que necessita tanto do canal visual quanto auditivo para que tenha efeito completo. Dependendo do tipo de legenda produzida, pode ser considerada como uma forma de tradução

intralinguística, interlinguística ou intersemiótica. É comum que essas categorias de legendas apresentem narrativas de natureza imagética, conhecidas dentro do mercado de legendagem como "on-screen texts" ou textos em tela. Nessa plataformas, em especial, o uso de textos em tela é bastante comum e, muitas vezes, competem com a informação de legenda que é apresentada. Muitas vezes, há a sobreposição de legendas com textos em tela, levando a uma redundância da informação. Outras vezes, produtores optam por retirar as legendas em trechos do vídeo onde o texto em tela apresenta a mesma informação textual. Sendo assim, o presente projeto pretende investigar o movimento ocular durante a exibição de legendas totais e parciais de vídeos com texto em tela. Os dados serão analisados com base em Modelos Lineares Generalizados Mistos ou Equações de Estimativas Generalizadas. Será realizado um estudo de metodologia descritiva qualitativa.

#### Objetivo da Pesquisa:

**Objetivo Primário:**

Investigar o movimento ocular durante a exibição de legendas totais e parciais de vídeos com texto em tela.

**Endereço:** Rua Cel. Nunes de Melo, 1000

**Bairro:** Rodolfo Teófilo

**UF:** CE

**Município:** FORTALEZA

**Telefone:** (85)3366-8344

**CEP:** 60.430-275

**E-mail:** comepe@ufc.br