

Publication status: This preprint has not been published elsewhere.

How comprehensible are eye-movement modelling examples? What annotations can tell

Juliana do Amaral, Davi Alves Oliveira, Mailce Borges Mota

<https://doi.org/10.1590/SciELOPreprints.15902>

Submitted on: 2026-04-21

Posted on: 2026-09-18 (version 2)
(YYYY-MM-DD)

Version justification: Versão submetida à revista Mental Lexicon.

How comprehensible are eye-movement modelling examples? What annotations can tell

Juliana do Amaral

Federal University of Santa Catarina

ORCID <https://orcid.org/0000-0003-1024-6066>

Davi Alves Oliveira

State University of Bahia

ORCID <https://orcid.org/0000-0001-6834-0579>

Mailce Borges Mota

Federal University of Santa Catarina

ORCID <https://orcid.org/0000-0002-8674-2480>

Author Note

We gratefully acknowledge funding from CNPq public notice nº 32/2023 for junior postdoctoral fellowships. Correspondence concerning this article should be addressed to profjulianoamaral@gmail.com.

Abstract

Eye movement modelling examples (EMMEs) are a type of video-based instruction that uses a model's gaze to draw the viewer's attention directly to the procedures that are being taught. Its uses extend across areas such as sports, aviation, medical education and reading instruction. The assumption is that watching a model may prompt the viewer to apply these procedures in similar tasks in the future, and visualizing good and poor behavior can trigger contrasting analysis. While the effects of EMMEs are usually measured in transfer tasks and model-observer similarity, few studies employed online comprehension checks. In the present study, we analyze the annotation protocol developed by do Amaral et al. (2025), who used EMMEs to foster digital reading competencies. The notes consisted of model descriptions and justified ratings. We hypothesized that annotations were not only a comprehension check but also a metacognitive tool that would foster processing of the EMMEs. We also investigated a possible correlation between annotation scores and performance in the transfer tasks of the previous study. Results indicated that participants understood the video models, as evidenced by high scores in the descriptions and ratings. However, qualitative analyses of the justifications revealed that participants attended to content features that were not fixated by the model. Content salience is justified in light of the purpose of the EMMEs, which aimed at developing digital reading competencies. Annotation scores did not correlate with transfer measures. We interpret this result as indicative of the discreteness between metacognition and observational learning, with the effect of EMMEs standing independently from written elaboration processes.

Keywords: EMMEs, annotation, comprehension check, digital reading competencies.

How comprehensible are eye-movement modelling examples? What annotations can tell

1. Introduction

The eye-tracking technology allows researchers to follow the eye movements of people as they process a visual stimulus. It has been used in multiple interdisciplinary fields such as neuroscience and psychology (including reading, scene perception, and visual search), engineering (exploring human-related factors in aviation and driving), marketing (e.g., visual inspection of advertisements, web pages, and labels), and computer science e.g., human-computer interaction (Duchowski, 2007). In reading, eye trackers are an accurate measure of allocation of visual attention. They provide data on whether, for how long, and in which order the viewer attends to each piece of information (Scheiter; Van Gog, 2009). Besides, this tool sheds light on cognitive processes that are unconscious to the reader and thus unlikely to be verbalized (in verbal protocols, for instance) nor evidenced in offline comprehension measures such as essays or comprehension tests.

Eye trackers capture a variety of eye data. As we read, our eyes sometimes stand still (fixations), or move between words or chunks (saccades). We also make backward movements to previously read words (regressions), and return sweeps at the end of the line to the next one below (which are often inaccurate). In investigating the role of speed on comprehension, Just and Carpenter (1980) analysed the gaze duration of college students and proposed the *eye-mind assumption*, according to which the pauses that readers make were indices of increased processing load. This load was caused by identification of low-frequency words, integration between important clauses, and generation of inferences at the end of sentences. The *eye-mind assumption* served as the theoretical basis for subsequent studies investigating eye movements in reading comprehension.

The analysis of eye-tracking data varies according to the type of stimulus. When analyzing word processing, common eye movement measures used are first-fixation duration on a word, single-fixation,

gaze duration, and total fixation time (Rayner et al., 2006). Analyses at the sentence or discourse level commonly measure first pass reading time and total reading time (Rayner, et al., 2005). Although longer reading time is usually interpreted as comprehension difficulty, it can also be a sign of increased attention and strategic processing. Table 1 summarizes the common terms and the definitions of common measures used in eye-tracking data analysis.

Table 1*Common Terms and Measures in Eye Tracking Research*

Term	Definition
Fixation	A period of relative ocular stability during which the gaze remains focused on a specific location, lasting roughly 200-250ms. ^a
Saccade	Rapid eye movements that shift the fovea between points of relative stability. In reading, saccades comprise seven to nine letter spaces. ^a
Regression	A movement directed to a segment of text or a stimulus area that has already been processed. Regressions are usually interpreted as index of comprehension difficulty (Rayner et al., 2005)
First-fixation duration	The duration of the very first fixation on a target word or area of interest.
Single-fixation duration	The duration of the fixation in instances where a target area receives only one fixation before the gaze moves to a different target.
Gaze duration	The sum of all fixation durations on a target during the first encounter, before the gaze leaves that area for the first time.
Total fixation time	The sum of all fixations made on a target (e.g., single word) prior to moving to the next one.
First pass reading time	The sum of fixations on a region (e.g., group of words) before moving forward.

Total reading time	The sum of all fixations in a region (e.g., group of words) or area of interest.
--------------------	--

^a These figures vary depending on the comprehension problems faced by the reader (Rayner, 1998).

Eye-tracking data can be visualized as statistical output and as video recording. The eye-movement recordings comprise the person's gaze superimposed to the stimulus. In addition to the investigation of visual processing, these gaze recordings have been used to investigate sentence processing and monitor strategic reading online (Hautala et al., 2018; Santos, 2019). Gaze recordings are the eye-tracking output used to create eye-movement modelling examples (henceforth EMMEs). EMMEs are selected parts of gaze recordings of a model performing a visual task. They are often research-generated to ensure clear modelling. The assumption is that, by watching a model perform a task at an optimal level, the viewer may apply these procedures in similar tasks in the future. The effectiveness of EMMEs on attention guidance and task performance is well established (Xie et al., 2021; Emhardt et al., 2023). However, it remains unclear how this video model is understood by viewers – which is the gap this study aims to explore.

In investigating the effect of EMMEs on the development of digital reading competencies in L2, Do Amaral et al. (2025) included an annotation protocol. It consisted of a request for descriptions and ratings of the models, to be answered at predetermined pauses as participants watched the video. These pauses for annotation are believed to induce deeper processing and reduce cognitive overload (Aalioui et al., 2022). In this paper, we analyse the notes and ratings done by participants as they watched the video. The assumption is that these notes can unveil whether EMMEs were correctly interpreted, any difficulties faced by participants in trying to understand them, and how accurately the video was rated.

1.1 Eye-movement modeling examples (EMMEs)

Eye-movement modeling examples are a type of example-based instruction in which a model shows how to do a task or to process visual material (e.g., a diagram). However, while in traditional example-based videos the model is visible as they perform the task, eye-movement modelling shows the recording of a task (often performed onscreen) from the perspective of the model, with the model's gaze displayed over the stimuli (Van Gog et al., 2009). The viewer's attention can be guided by showing a moving dot (representing the model's exact fixation), heat maps (with hot shades indicating the areas that were fixated for longer), or spotlight which blurs the areas that should be ignored (Emhardt et al., 2023). EMMEs usually display excellent performance, but low-performing models can also be used to trigger contrastive analysis (Salmerón & Llorens, 2019; Krebs et al., 2021; Do Amaral et al., 2025). These videos are considered to be more effective than regular video modelling for instructional learning because the viewer's focus of attention synchronizes with the model's (Chisari et al., 2020).

This video modelling technique is employed for different goals depending on the area of expertise. For instance, in sports and aviation, their use is associated with strategy development, while in the STEM areas they are used to guide attention (Emhardt et al., 2023). In business education, EMMEs were used to enhance comprehension of process models i.e., graphic representations of all the steps involved in services such as cash and order management (Winter et al., 2021). In medical education, they were employed to teach medical students what to observe during training on clinical diagnosis (Jarodzka et al., 2012) and physiology lectures (Aalioui et al., 2022). In educational psychology, EMMEs have been explored in their potential to draw attention to the visual features that are assessed in a subsequent perceptual task (Jarodzka et al., 2013), or to teach text-picture integration to seventh-graders (Mason et al., 2016) and university students (Scheiter et al., 2018).

Recent meta-analyses (Xie et al., 2021; Emhardt et al., 2023) found positive effects of EMMEs on increased attention to the areas modelled, and on performance in classification tasks, visuo-monitor

tasks, and text processing and comprehension. However, it is less clear whether EMMEs increase performance on problem-solving and learning tasks (van Marlen et al., 2016; Emhardt et al., 2023). The effect of EMMEs on learning is dependent on prior knowledge. Low prior knowledge learners tend to benefit more from the intervention compared to learners with high prior knowledge on the topic (Krebs et al., 2019; Chisari et al., 2020, Aalioui et al., 2022), although studies show non-converging results (i.e., Scheiter et al., 2018; Krebs et al., 2021).

EMMEs can either be silent or narrated, including the comments of the model and/or a verbal description of the steps (Emhardt et al., 2023). Silent EMMEs isolate the gaze component of instruction (Mason et al., 2015) – which facilitates control over the predictor variable. On the other hand, using EMMEs with little to no verbal (written and/or oral) explanation may increase extraneous cognitive overload (Sweller et al., 2011) since the learner has to make more inferences to understand what is being modelled. Therefore, manipulation checks such as post task interviews and annotations arise as important measures when implementing this type of video-based instruction. They are approached in the next section.

1.2.1 Measuring the effectiveness of EMMEs

Studies investigating the effectiveness of EMMEs tend to compare one or more experimental groups to a control group which either receives no treatment (Mason et al., 2015, 2016; Krebs et al., 2021), or a similar video-based condition (Salmerón et al., 2020; Do Amaral et al., 2025). The design of the control condition plays an important role in assessing the effectiveness of a study using EMMEs. Salmerón and colleagues (2020) used a video about the content of the navigation task participants would perform subsequently. Do Amaral, Oliveira and Salmerón (2025) adopted a more instruction-driven approach: the control group watched a video with explicit instruction on online reading strategies.

Previous studies investigated the effectiveness of EMMEs using pre/posttests – including recall and transfer tasks (Salmerón et al., 2020; Mason et al., 2015; 2017; Chisari et al., 2020), and scan path similarity between the viewer's gaze and the model's (Jarodzka et al., 2012; 2013; Krebs et al., 2019). However, even accurate measures such as the Euclidean distance (Jarodzka et al., 2012) do not unveil participants' reasoning processes as they watch the EMMEs.

Mason and colleagues (2015, 2016) did a post task interview as manipulation task to ensure participants understood the behavior that was modelled by the eye movements (i.e., text-picture integration when reading an illustrated text) and the aim of the video (i.e., to model text-picture processing). To our knowledge, these were the two first and only EMME studies to implement a manipulation check. However, the data were collected retrospectively, that is, after participants had watched the EMMEs. Online manipulation checks such as annotations as participants watch the video have not yet been explored in their potential to clarify the processing of EMMEs.

1.2 The case for annotations

When watching a video, people with normal vision and hearing process information automatically through their visual and auditory channels. However, learning situations require more active engagement with the material. Verbally elaborating about the behavior observed is believed to enhance retention and transfer (Bandura, 1986). In reading, study strategies such as underlining/highlighting, annotating, and self-generated questions help the reader organize textual information in a way that facilitates retrieval (Just & Carpenter, 1980; Nist & Simpson, 2000). Its benefits extend to learning from L2 texts (Do Amaral & Tomitch, 2022) and to audio and video (Baron, 2021).

Note taking is an active study strategy because it promotes cognitive engagement and deep processing by (re)organizing the information through elaboration processes (Craik & Lockhart, 1972). In addition, this strategy has *external storage effects*, since the notes that were produced become study material (Jansen et al., 2017). Taking notes also yields more time for reasoning compared to verbal

protocols and interviews. The modality in which the content is presented (text or video) does not seem to influence encoding – however existing evidence is scarce and lacks statistical power (Jansen et al., 2017). In relation to how the notes are taken, longhand note-taking does not lead to better comprehension compared to laptop typing (Urry et al., 2021). It is rather the depth of processing that determines the quality of the notes and their encoding effects.

In lectures and video-based tasks, note-taking decreases mind wandering by helping learners keep focus, enhances cognitive engagement, and can be used to organize information for later retrieval in the form of list of keywords, summaries, or concept maps (Fang et al., 2022; Baron, 2021). However, because the learners take notes as they listen, the resulting notes are more often transcriptions than elaborations, with positive effects of transcribed notes on memory (Bui et al., 2013). Taking notes also fosters sustained attention when performing an online search task under time pressure (Wu & Xie, 2018). The development of the note taking strategy requires time, effort, and training (Rawson & Kintsch, 2005). Among undergraduate students and majors, metacognitive instruction on how to take notes from oral input (i.e., lectures and videos) led to positive results on summary writing and memory for information units in L1 and L2 (see Siegel 2021 for an overview). However, the majority of the studies investigating note taking focus on its effectiveness on memory. Studies using notes as a metacognitive tool to check comprehension of video-based instruction – which is the goal of the present study – are more scarce.

1.3 The present study

The goal of this study is to gain a more accurate understanding of how EMMEs contribute to model behaviour in general, and digital reading competences in particular. To this aim, we analysed the annotation protocol, which asked participants to provide a) a description of the behaviour modelled and b) a justified rating of each model on a 5-point Likert scale. The descriptions were analysed using a rubric (Do Amaral et al., 2025), and the ratings were coded. Justifications to the ratings were analysed

qualitatively. We argue that annotations provide an accurate picture of learner's processing of eye movement models, since it allows self-monitoring and elaboration about the strategies depicted in the video. Two research questions (RQ) were posed:

1. Do annotations and Likert scales evidence accurate description and justified rating of the competencies modelled by the EMMEs?
2. Do annotation composite scores (i.e., description and rating) correlate with performance in the transfer tasks i.e., navigation and evaluation of webpage reliability?

Do Amaral, Oliveira and Salmerón (2025) found transfer effects of instruction with EMMEs on transfer tasks measuring navigation (evidenced by increased time inspecting a search engine results page – SERP) and evaluation (evidenced by decreased time on non-reliable webpages). However, no effects of group were found on evaluation of source features. Therefore, in RQ1 we hypothesize that annotation scores have a positive correlation with navigation and evaluation of webpage reliability. Source evaluation was not included since no effects of EMMEs on fixation to source features were observed in the previous study.

In relation to RQ2, we also hypothesize a positive correlation between higher scores in the annotation protocol and the processing measures analysed in the previous study. Additionally, EMMEs tend to focus on excelling performance as the standard instructional approach. Therefore, we hypothesized that good models would be better understood and more accurately rated and justified compared to poor models because good models provide more information, while poor models require that the viewer makes more inferences to interpret the video (e.g. by inferring wrong procedures and missing steps in the process).

2. Method

2.1 Participants and design

The data consists of a subsample from a previous study (Do Amaral et al., 2025) which investigated the role of EMMEs in modelling digital reading competencies. Participants were 29 European undergraduate students speakers of English as L2 (25 female, M age = 20,86, SD = 3.35). The subsample selected consisted of the participants in the experimental (EMME) condition, i.e., the ones who watched the video models and took notes.

2.2 Materials

2.2.1 Eye-movement modelling examples (EMMEs)

The EMMEs developed by Do Amaral and colleagues (2025) was a 9'32'' video that showed the gaze path of eight models as they performed navigation tasks. The video was divided into two sections. The first section showed models who did well in the task, while the second showed learners who did not perform so well. This instructional approach derived from perceptual learning theory is known as contrasting cases (Schwartz & Bransford, 1998). It fosters analysis with a focus on differentiation and can be applied for images or textual material. Analyzing contrasting cases is believed to foster metacognition and elaboration. In using EMMEs with contrasting cases to model digital reading behavior, Salmerón and Llorens (2019) found a positive effects of EMMEs on digital text comprehension.

The strategies modelled by the good models comprised 1) SERP (search engine results page) inspection, 2) identification of source information, 3) deep reading of trustworthy and relevant pages, 4) skimming of less trustworthy pages, and 5) quickly abandoning topically unrelated pages. Models 6, 7 and 8 (who did not perform well) failed to apply strategies 1) SERP inspection, 2) identification of source information, and 3) deep reading of trustworthy and relevant pages (see appendix B for a full description of the strategies modelled). As forestated, contrasting good and poor behavior has the potential to foster elaboration. However, because the strategies were displayed as eye movements,

poor performance was represented as little visual inspection – which could pose a greater challenge to participants in understanding this type of modelling. Therefore, the unequal balance between good and poor models is justified for two reasons. First, it exposes the participant more to optimal than to incorrect behaviour. Second, it enables comparison of how participants comprehend good and poor models performing the same strategy.

2.2.2 Annotation protocol

The annotation protocol in which participants assessed EMMEs was the online manipulation check developed by Do Amaral et al. (2025). It aimed to investigate whether the competencies modelled were understood by participants, and the accuracy of their ratings and justifications. The worksheet comprised a) a description of model behaviour and b) rating of the models (five-point Likert scale ranging from “very poor” to “very good”) followed by a justification. The annotations were taken on a worksheet provided by the researchers (Appendix A). For each of the eight models watched, participants answered two questions: “1) What does this student do?” and “2) On a 5-point scale, how do you rate this students' reading and analysis of the material? Justify your answer.” The Likert scale ranged from “very poor” to “very good”. The description of each model and the justification were coded as a single written response. Following the structure of the EMMEs, the worksheet was explicitly divided into two sections: “students who perform well” and “students who do not perform so well”. At the end of each section, participants were instructed to pause the video and take notes.

2.3 Data analysis

The annotations were analysed in relation to a rubric (Appendix B). Scores for the description of each of the five strategies modelled were 0 (no strategy mentioned), 1 (one strategy mentioned), or 2 (two or more strategies mentioned). Data that was illegible and/or erased was excluded from the sample. Data from the Likert scale ratings was coded differently for the models of good and poor performance. For the models of good performance, high scores in the Likert scale were interpreted as

better comprehension of the EMMEs (very poor = 1, poor = 2, average = 3, good = 4, very good = 5). Differently, for the models of poor performance, scores were reverse coded with high scores interpreted as inaccurate rating (very poor = 5, poor = 4, average = 3, good = 2, very good = 1). This procedure allowed finer-grained comparisons between comprehension of poor and good models and the computing of an annotation composite score. Thus, the three variables analysed were the sum of the ratings in the Likert scale for each participant (henceforth, rating scores) with possible values ranging from 0 to 40, the mean of the sum of the scores given by two raters for the description for each participant (henceforth, description scores) with possible values ranging from 0 to 16 , and the composite score, which is the sum of ratings and description, with possible values ranging from 0 to 56. Last, the justifications provided to the ratings were analysed qualitatively through the search of patterns for justification (Appendix C). Data is openly available at <https://doi.org/10.5281/zenodo.19680777>.

2.4 Use of AI-based tools

The authors incorporated OpenAI's ChatGPT (GPT-5.6 Sol), into their writing workflow for grammatical and orthographic review. Following this assistance, the authors performed thorough quality control and content revision as needed, maintaining full responsibility for all aspects of the published article.

3. Results

3.1 Preliminary analysis

The descriptive and statistical tests were done using JASP (2025) software for statistical analyses version 0.19.3. Descriptive data is presented in Table 2. We run a Spearman's correlation test to identify whether these two measures were correlated and, if so, the direction of the correlation. Spearman's correlation coefficient showed a significant positive relationship between scores in the annotations and Likert scale ratings ($\rho = 0.61$, $p < .001$), which endorses the internal consistency of the instrument.

Table 2

Descriptive statistics

	Model rating	Annotations
Valid	29	29
Mean	27.72	10.69
Std. Deviation	4.69	3.01
Skewness	-0.41	0.15
Std. Error of Skewness	0.43	0.43
Kurtosis	0.69	-1.09
Std. Error of Kurtosis	0.84	0.84
Shapiro-Wilk	0.97	0.94
Minimum	15.00	5.00
Maximum	37.00	15.50

3.2 Annotations: description and justified rating of the models

The descriptions were scored by the first and second authors and we used the mean of the sum of scores for each participant as the description variable. The correlation between scores of the first and second raters was high ($r = 0.880$, $R^2 = 0.774$). Scores were normally distributed ($W = 0.94$) and mean score in the annotations was 10.69 ($min = 5.00$, $max = 15.50$, $SD = 3.01$), which is evidence of overall good performance. To answer RQ1 “Do annotations and Likert scales evidence accurate description and justified rating of the competencies modelled by the EMMEs?”, we analysed the amount of strategies mentioned in the descriptions and how well participants rated models of good performance as good and models or poor performance as poor in the ratings.

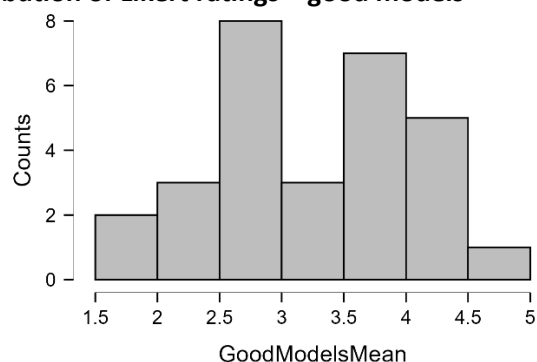
For description scores, we considered that participants should mention at least one strategy. A case in which only one strategy was mentioned for all questions would result in a score of 8, and a case in which two or more strategies were mentioned for all questions would result in a score of 16. Thus, we expected the average of the description scores to be between 8 and 16 and we used the mean of these two values, 12, as the expected values in a one sample T-test. The test showed a statistically significant

difference between the mean of description scores ($M = 10.69$, $SD = 3.01$) and the expected value, 12 ($t(28) = -2.34$, $p = .026$), which indicates that our expectation was not fulfilled.

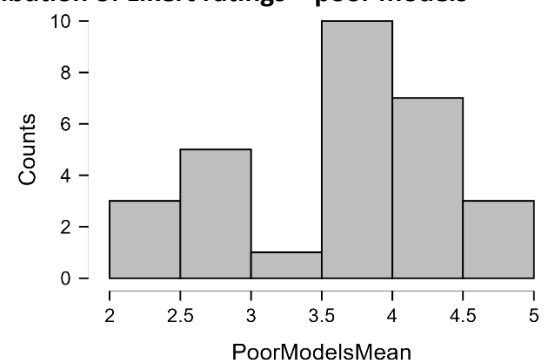
For rating scores, we considered that models of good performance should be rated good (4 points in the Likert scale) or very good (5 points in the Likert scale) and that models of poor performance should be rated poor (4 points after reverse coding) or very poor (5 points after reverse coding). Considering that each participant rated 8 models, we expected the average of the rating scores to be between 32, which represents a case in which all answers were 4, and 40, which represents a case in which all answers were 5. Thus, we used as expected value the mean of these two values, 36, and we ran a one sample T-Test to test for a statistically significant deviation of this value. The test showed a statistically significant difference between the mean of the sum of the ratings ($M = 27.72$, $SD = 4.70$) and the expected value, 35 ($t = -9.49$, $p < .001$), which indicates that our expectation was not fulfilled. Median values for the rating of the good models were 3.40, while for the poor models median was 4.0, which indicates that poor models were slightly better rated than their positive counterparts. The distribution plots (Figure 1) also illustrate a more consistent distribution of high scores among poor models.

Figure 1 – distribution plots

Distribution of Likert ratings – good models



Distribution of Likert ratings – poor models



The analysis of the notes reveal that participants were able to correctly infer the strategies that were modelled, as evidenced by high scores in the descriptions, although not as high as expected. However, a qualitative analysis of the justifications provided for the ratings indicate that analysis of model's behaviour was based not only on the gaze but also on other features of the webpage that were not fixated by the model. These features were most often related to the type and content of the webpage. For instance, in strategy 2 (identification of source information), the model enters a webpage, looks at the logo (a mineral water brand), reads the text and carefully inspects the author information provided below the text. Although this was a good model of source evaluation, it was perceived as poor behavior by eight participants (appendix C) because the webpage was commercial. Therefore, participants concluded that it should not have been visited.

Similarly, in analysing the model of strategy 4 (skimming of less trustworthy pages) nine participants perceived access to the unreliable webpage (i.e., popular forum) as poor performance (appendix D), ignoring the fact that it was only skimmed. Curiously, among those, there were four mentions of excessive attention to the source features – which became a salient gaze behaviour considering that the webpage was not a credible source and the person who posted the content was a layperson. In strategy 5 (quickly abandoning a topic-unrelated page), quick access to an unrelated webpage (i.e., a waterpark) was perceived as poor behaviour by nine participants. Taken together, these results suggest that participants failed to notice that accessing less trustworthy and task-unrelated webpages was necessary to model evaluation of its source features, the reliability of its content, and pertinence to the task – which resulted in accurate descriptions but low ratings of the EMMs. In addition, type of webpage and content were as salient (or even more) than the model's gaze point – particularly in the analysis of strategies 2, 4 and 5. We further debate content salience in source evaluation in the discussion.

The skimming strategy was misinterpreted by 15 participants when watching models 4 (skimming of less trustworthy pages) and 5 (quickly abandoning topically unrelated pages). We attribute this result to a spillover effect of strategy 3 (deep reading of trustworthy and relevant pages) over the analysis of the subsequent strategies modelled. To be precise, previous view of strategy 3 might have affected participants' subsequent model analysis by setting a "read all" standard – which in fact should not be a criterion for the analysis of the models evaluating the unreliable webpages (strategies 4 and 5). Reading all the content of the webpage was then generalized as good behaviour – and models who haven't done so were rated as poor.

3.3 Annotation composites and transfer tasks

To answer RQ2 "Do annotation composite scores (i.e., description and rating) correlate with performance in the transfer tasks i.e., navigation and evaluation of webpage reliability?", we ran three linear models. First, a correlation was analysed between the variables total fixation time per word on the search engine results page (SERP) and annotation composite scores. The correlation was not statistically significant ($r = 0.10$, $p = 0.57$, 95% CI = [-0.26, 0.45]). Second, we analysed another correlation with fixation time per word on reliable pages and annotation composite scores. The correlation was also not statistically significant ($r = 0.22$, $p = 0.23$, 95% CI = [-0.54, 0.15]). Last, we analysed a correlation with total fixation duration per word on source features and annotation composite scores, but again the correlation was not statistically significant ($r = 0.08$, $p = 0.66$, 95% CI = [-0.29, 0.43]). Therefore, no correlations were found between the annotations written while watching the EMMs and performance in the transfer tasks.

4. Discussion

The present study aimed to examine the processing of eye movement modelling examples by analyzing the annotations taken by participants as they watched a video aimed at modelling digital

reading competencies such as navigation across search engine results and evaluation of the source characteristics and the content reliability of the resulting webpages (Do Amaral et al., 2025). The assumption was that these notes inform how accurately participants understood and rated the behavior modelled in the video, and the quality of the justifications provided. The insights derived from this analysis also inform enhancements in eye movement modelling in general and particularly in instruction on digital reading competencies.

4.1 Model ratings

The EMMEs were divided into two sections. The first showed eye movements of people who performed well in the task, whereas the second comprised participants who did not perform so well. To make clear this distinction, written instructions were provided both in the video and in the annotation worksheet. However, as represented in the distributions plots (Figure 1), some participants rated good performances as “poor” and/or “very poor” – which is at odds with our initial assumption that good models are easier to understand while poor models require greater inference generation. One possible explanation is the fact that participants had been previously exposed to the good model performing the same strategy, which facilitated analysis under the contrasting cases paradigm (Schwartz & Bransford, 1998). Therefore, our result corroborates Salmerón and Llorens (2019) who found positive effects of video modelling with contrasting cases on the development of digital reading competencies.

4.2 Content as a sourcing criterion

Participants justified their ratings based on content characteristics of the webpages rather than on gaze behaviour i.e., how model attention was being allocated to specific portions of the webpage. For instance, in analysing the model of strategy 2 (identification of source information) participants rated access to a commercial webpage (mineral water) as poor behaviour. This result indicates that the textual content in which modelling is embedded (i.e., the website) bears great salience to the viewer – even when they are explicitly asked to focus on the model’s gaze.

The salience of content in the EMMEs developed for this study can be explained in terms of the type of behaviour modelled: digital competencies in multiple document reading (Salmerón et al., 2018). Strategies such as inspecting the search engine results, identifying source information, deep reading of trustworthy pages, skimming of less trustworthy pages, and quickly abandoning topically unrelated pages are deeply rooted in content analysis. Content analysis is emphasized in theoretical frameworks of sourcing and multiple document comprehension such as the Documents Model (Perfetti et al., 1999), and the Content-source integration (CSI) model (Stadtler & Bromme, 2014). In the Documents Model, a mental representation of the texts is constructed by the reader comprising text information such as source features (author, setting and form), content (main ideas), and rhetorical goals, as well as the relations among texts and from the texts to their content. When the documents present opposing perspectives, the CSI model poses that readers attempt to restore coherence (by reconciling the perspectives or attributing the conflict to different views) and solve the conflict by judging the validity of the arguments and the quality of the source.

In sum, both theoretical models assume that content analysis is an important component of source evaluation. However, the issue of how to balance content and gaze salience in EMMEs remains to be investigated. Future studies should acknowledge content salience and further manipulate its potential by avoiding the use of blurred video tools and by adapting EMMEs to specific topics and languages. Because eye movements are at the core of EMMEs, it is also crucial to increase their salience. One alternative is to reinforce the viewer's attention to the "moving dot" in the video instructions. Another is to adjust the time window of the gaze tail to display more gaze plots. The annotation sheets could also drive the viewer's focus to the gaze through instructions such as "please follow the model's gaze. Where does she fixate attention?" to prevent participants from attending to text that is not signalled by the model's gaze.

4.3 Observational learning and metacognition

In the previous study, we found a statistically significant effect of group on navigation and evaluation: participants who watched the EMMEs performed better on a subsequent navigation task while their eye movements were tracked. They spent more time inspecting the search engine results and less time on non-reliable webpages compared to controls. The results were interpreted as effect of instruction on digital reading competencies with EMMEs. However, we did not find group differences in terms of total fixation on the source features i.e., participants in the experimental group (who watched the EMMEs) performed similarly to controls when inspecting the authorship of the webpages (Do Amaral et al., 2025).

In the present study, we expected that annotations would positively correlate with navigation and evaluation performance – which was not confirmed. This can be explained in terms of metacognitive overload (Sweller, 2011). Participants were able to reproduce the strategies they were taught in the EMMEs in a subsequent navigation task. However, they could not elaborate about them as they watched the instructional video. Although this is at odds with seminal theories of observational learning (Bandura, 1986), the results suggest that the positive effects of EMMEs are independent from metacognitive processes of elaboration. In other words, watching the model's gaze was enough to form a conceptual representation of the reading behaviour observed and transfer it to a new task. Nevertheless, confirmation checks remain a relevant methodological procedure. Because EMMEs do not convey explicit verbal instruction, they require a great deal of inference making from the viewer. More straightforward confirmation checks such as open-ended questions or retrospective think aloud protocols could be used instead of annotations to ensure learners understood what was modelled.

EMMEs also need to be further explored in their potential to inform instruction under more ecologically valid contexts such as basic education. In this context, metacognitive elaboration can be triggered under an instructor-controlled pace to provide some scaffolding and ensure comprehension. However, the processing challenge involved in watching a model, understanding what is being shown

and why this is considered optimal (or poor) behaviour is a welcome effort that should be preserved when providing students with metacognitive scaffolding during instruction with EMMEs.

5. Conclusion

Eye movement modelling examples (EMMEs) are a complex balance between content elements (such as text, image or video) and gaze information that may be well comprehended by the viewer and successfully transferred to new learning situations. However, comprehension measures need to be added to instruction to make sure participants understand what is modelled – and including annotations might have led to some cognitive overload. Although participants in the present study provided accurate descriptions and appropriate ratings of the behaviour modelled, mean behaviour deviated from our expectations. Taken together, the analysis of annotations provides accurate insights into online comprehension of this type of instruction – and particularly on how textual content acts as a criterion when watching EMMEs, diverting the viewer's attention from the model's gaze. Future studies could explore more subtle scenarios of evaluation in which content plays a prominent role. For instance, textual sources written by similarly qualified authors with divergent views on the same topic would trigger deep content analysis – and modelling the strategies to be used in these contexts could help readers acknowledge how complex it is to decide what to believe in a world of misinformation.

Authors contributions (*contribuição de autoria*)

Juliana do Amaral: conceptualization, data curation, formal analysis, funding acquisition, investigation, methodology, project administration, writing – original draft preparation, review and editing. **Davi Alves Oliveira:** data curation, formal analysis, investigation, methodology, writing – original draft preparation, review and editing. **Mailce Borges Mota:** data curation, formal analysis, funding acquisition, investigation, project administration, supervision, writing – review and editing.

Conflicts of interest (*conflito de interesses*)

We have no conflicts of interest to disclose.

Ethics committee approval (*declaração de aprovação do Comitê de Ética*)

The project was approved by the ethics committee of University of Valencia, Spain (register nº 1822514).

Data availability statement (*declaração de disponibilidade de dados de pesquisa*)

Data is openly available at <https://doi.org/10.5281/zenodo.19680777>.

AI use statement (*declaração de uso de inteligência artificial*)

During the preparation of this manuscript, the authors did not use any Artificial Intelligence tools.

References

Aalioui, L., Gouzi, F., & Tricot, A. (2022). Reducing cognitive load during video lectures in physiology with eye movement modeling and pauses: a randomized controlled study. *Advances in physiology education*, 46(2), 288–296. <https://doi.org/10.1152/advan.00185.2021>

Bandura, A. (1986). Observational Learning. In A. Bandura, *Social foundations of thought and action* (pp. 47-105). Englewood Cliffs, NJ.

Baron, N. S. (2021). *How we read now: strategic choices for print, screen, and audio*. New York: Oxford University Press. <https://doi.org/10.1093/oso/9780190084097.001.0001>

Brysbaert, M. (2019). How many words do we read per minute? A review and meta-analysis of reading rate. *Journal of Memory and Language*, 109, 1-30. <https://doi.org/10.1016/j.jml.2019.104047>

Bui, D. C., Myerson, J., & Hale, S. (2013). Note-taking with Computers: Exploring alternative strategies for improved recall. *Journal of Educational Psychology*, 105(2), 299–309. <https://doi.org/10.1037/a0030367>.

Chisari, L. B., Mockevičiūtė, A., Ruitenburg, S. K., van Vemde, L., & Kok, E. M. (2020). Effects of prior knowledge and joint attention on learning from eye movement modeling examples. *Journal of Computer Assisted Learning*. First published: 27 February 2020 <https://doi.org/10.1111/jcal.12428>.

Craik, F. I. M., & Lockhart, R. S. (1972). Levels of processing: a framework for memory research. *Journal of Verbal Learning and Verbal Behavior*, 11(6), 671-684. [https://doi.org/10.1016/S0022-5371\(72\)80001-X](https://doi.org/10.1016/S0022-5371(72)80001-X)

Do Amaral, J., Salmerón, L., & Oliveira, D. A. (2025). Eye-Movement Modelling as a Tool to Foster Digital Reading Competencies and Misconception Change When Reading Webpages in L2. *Journal of Computer Assisted Learning*, 41(4), e70091. <https://doi.org/10.1111/jcal.70091>.

Do Amaral, J., & Tomitch, L. M. B. (2022). Comparing the effectiveness of study strategies on comprehension, retention, and learning from L2 English texts. *Letras de hoje*, 57(1), e41961-e41961. <https://doi.org/10.15448/1984-7726.2022.1.41961>.

Duchowski, A. T., & Duchowski, A. T. (2017). *Eye tracking methodology: Theory and practice*. Springer. <https://link.springer.com/book/10.1007/978-3-319-57883-5>

Emhardt, S.N., Kok, E., van Gog, T. *et al.* (2023). Visualizing a Task Performer's Gaze to Foster Observers' Performance and Learning – a Systematic Literature Review on Eye Movement Modeling Examples. *Educ Psychol Rev* 35, 23. <https://doi.org/10.1007/s10648-023-09731-7>

Fang, J., Wang, Y., Yang, C. L., Liu, C., & Wang, H. C. (2022). Understanding the effects of structured note-taking systems for video-based learners in individual and social learning contexts. *Proceedings of the ACM on Human-Computer Interaction*, 6, 1-21. <https://doi.org.ez46.periodicos.capes.gov.br/10.1145/3492840>

Hautala, J., Kiili, C., Kammerer, Y., Loberg, O., Hokkanen, S., & Leppänen, P. H. T. (2018). Sixth graders' evaluation strategies when reading Internet search results: an eye-tracking study. *Behaviour & Information Technology*, 37(8), 761–773. <https://doi.org/10.1080/0144929X.2018.1477992>

Jansen, R. S., Lakens, D., & IJsselsteijn, W. A. (2017). An integrative review of the cognitive costs and benefits of note-taking. *Educational Research Review*, 22, 223-233.

<http://dx.doi.org/10.1016/j.edurev.2017.10.001>

Jarodzka, H., Balslev, T., Holmqvist, K. *et al.* (2012). Conveying clinical reasoning based on visual observation via eye-movement modelling examples. *Instr Sci* 40, 813–827.

<https://doi.org/10.1007/s11251-012-9218-5>

Jarodzka, H., van Gog, T., Dorr, M., Scheiter, K., & Gerjets, P. (2013). Learning to see: Guiding students' attention via a Model's eye movements fosters. *Learning and Instruction* 25, 62-70.

<https://doi.org/10.1016/j.learninstruc.2012.11.004>

Just, M. A., & Carpenter, P. A. (1980). A theory of reading: from eye fixations to comprehension. *Psychological review*, 87(4), 329-354. <https://doi.org/10.1037/0033-295X.87.4.329>

Krebs, M-C., Schöler, A., & Scheiter, K. (2019). Just follow my eyes: The influence of model-observer similarity on Eye Movement Modeling Examples. *Learning and Instruction*, 61, 126-137.

<https://doi.org/10.1016/j.learninstruc.2018.10.005>

Krebs, M-C., Schöler, A. & Scheiter, K. (2021). Do prior knowledge, model-observer similarity and social comparison influence the effectiveness of eye movement modeling examples for supporting multimedia learning?. *Instr Sci* 49, 607–635. <https://doi.org/10.1007/s11251-021-09552-7>

Santos, S., Maia, P., Albuquerque, T., & Aquino, R. (2019). Educação e trabalho: concepção politécnica na formação de leitores. In M. Maia (Org.), *Psicolinguística e metacognição na escola* (1st ed., pp.43-74). Mercado de Letras.

Mason, L., Pluchino, P., & Tornatora, M. C. (2015). Eye-movement modelling of integrative reading of an illustrated text: effects on processing and learning. *Contemporary Educational Psychology*, 41, 172-187. <http://dx.doi.org/10.1016/j.cedpsych.2015.01.004>

Mason, L., Pluchino, P., & Tornatora, M. C. (2016). Using eye-tracking technology as an indirect instruction tool to improve text and picture processing and learning. *British Journal of Educational Technology* 47 (6), 1083-1095. <https://doi.org/10.1111/bjet.12271>

Mason, L., Scheiter, K., & Tornatora, M. C. (2017). Using eye movements to model the sequence of text–picture processing for multimedia comprehension. *Journal of Computer assisted learning*, 33, 443–460. <https://doi.org/10.1111/jcal.12191>

Nist, S. L., & Simpson, M. L. (2000). College studying. In M. L. Kamil, P. B. Mosenthal, P. D. Pearson, & R. Barr (Eds.), *Handbook of reading research* (Vol. 3, pp. 645–666). Mahwah, NJ: Erlbaum.

Perfetti, C. A., Rouet, J. F., & Britt, M. A. (1999). Toward a theory of documents representation. *The construction of mental representations during reading*. Lawrence Erlbaum Associates.

Rayner, K. (1998). Eye movements in reading and information processing: 20 years of research. *Psychological Bulletin*, 124(3), 372–422. <https://doi.org/10.1037/0033-2909.124.3.372>

Rayner, K., Chace, K. H., Slattery, T. J., & Ashby, J. (2006). Eye Movements as Reflections of Comprehension Processes in Reading. *Scientific Studies of Reading*, 10(3), 241–255. https://doi.org/10.1207/s1532799xssr1003_3

Rayner, K., Juhasz, B. J., & Pollatsek, A. (2005). Eye Movements During Reading. In M. J. Snowling & C. Hulme (Eds.), *The science of reading: A handbook* (pp. 79–97). Blackwell Publishing. <https://doi.org/10.1002/9780470757642.ch5>

Salmerón, L., & Llorens, A. (2019). Instruction of digital reading strategies based on eye-movements modeling examples. *Journal of Educational Computing Research*, 57(2), 343–359. <https://doi.org/10.1177/0735633117751605>

Salmerón, L., Strømsø, H. I., Kammerer, Y., Stadtler, M., & Van den Broek, P. (2018). Chapter 4. Comprehension processes in digital reading. In *Learning to read in a digital world* (pp. 91-120). John Benjamins Publishing Company. <https://doi.org/10.1075/swll.17>

Scheiter, K., Schubert, C., & Schüler, A. (2018). Self-regulated learning from illustrated text: Eye movement modeling to support use and regulation of cognitive processes during learning from multimedia. *Br J Educ Psychol*. 88(1), 80-94. <https://doi.org/10.1111/bjep.12175>

Scheiter, K., & Van Gog, T. (2009). Using eye tracking in applied research to study and stimulate the processing of information from multi-representational sources. *Applied Cognitive Psychology*, 23(9). <https://doi.org/10.1002/acp.1524>

Schwartz, D. L., & Bransford, J. D. (1998). A Time For Telling. *Cognition and Instruction*, 16(4), 475-5223. http://dx.doi.org/10.1207/s1532690xc1604_4

Siegel, J. (2021). Retentional effects of notetaking instruction in EAP courses. *ELT Journal*, 75(3), 311-319. <https://doi.org/10.1093/elt/ccaa067>

Stadtler, M., & Bromme, R. (2014). The content–source integration model: A taxonomic description of how readers comprehend conflicting scientific information. *Processing inaccurate information: Theoretical and applied perspectives from cognitive science and the educational sciences*, 379-402.

Sweller, J., Ayres, P., & Kalyuga, S. *Cognitive Load Theory*. (2011). New York: Springer. <https://doi.org/10.4300/JGME-5-4-18>

Urry, H. L., Crittle, C. S., Floerke, V. A., Leonard, M. Z., Perry III, C. S., Akdilek, N., ... & Zarrow, J. E. (2021). Don't ditch the laptop just yet: A direct replication of Mueller and Oppenheimer's (2014) study 1 plus mini meta-analyses across similar studies. *Psychological Science*, 32(3), 326-339. <https://doi.org/10.1177/0956797620965541>

Van Gog, T., Jarodzka, H., Scheiter, K., Gerjets, P., & Paas, F. (2009). Attention guidance during example study via the model's eye movements. *Computers in Human Behavior*, 25(3), 785-791.

<https://doi.org/10.1016/j.chb.2009.02.007>

Van Marlen, T., van Wermeskerken, M., Jarodzka, H., & van Gog, T. (2016). Showing a model's eye movements in examples does not improve learning of problem-solving tasks. *Computers in Human Behavior*, 65, 448-459. <https://doi.org/10.1016/j.chb.2016.08.041>

Winter, M., Pryss, R., Probst, T., & Reichert, M. (2021). Applying Eye Movement Modeling Examples to Guide Novices' Attention in the Comprehension of Process Models. *Brain Sci.* 11, 72.

<https://doi.org/10.3390/brainsci11010072>

Wu, J-Y., & Xie, C. (2018). Using time pressure and note-taking to prevent digital distraction behaviour and enhance online search performance: Perspectives from the load theory of attention and cognitive control. *Computers in Human Behavior*, 88, 244-254.

<https://doi.org/10.1016/j.chb.2018.07.008>

Xie, H., Zhao, T., Deng, S., Peng, J., Wang, F., & Zhou, Z. (2021). Using eye movement modeling examples to guide visual attention and foster cognitive performance: A meta-analysis. *Journal of Computer Assisted Learning*, 37(4), 1194–1206. <https://doi.org/10.1111/jcal.12568>

Appendix A

EMMEs worksheet

Participant nº:

While you watch the eye movements, take notes:

PART 1 – STUDENTS WHO PERFORM WELL

STUDENT 1

What does this student do?

On a 5-point scale, how do you rate this students' reading and analysis of the material? Justify your answer.

Very poor	Poor	Average	Good	Very good
1	2	3	4	5

STUDENT 2

What does this student do?

On a 5-point scale, how do you rate this students' reading and analysis of the material? Justify your answer.

Very poor	Poor	Average	Good	Very good
1	2	3	4	5

STUDENT 3

What does this student do?

On a 5-point scale, how do you rate this students' reading and analysis of the material? Justify your answer.

Very poor	Poor	Average	Good	Very good
1	2	3	4	5

STUDENT 4

What does this student do?

On a 5-point scale, how do you rate this students' reading and analysis of the material? Justify your answer.

Very poor	Poor	Average	Good	Very good
1	2	3	4	5

STUDENT 5

What does this student do?

On a 5-point scale, how do you rate this students' reading and analysis of the material? Justify your answer.

Very poor	Poor	Average	Good	Very good
1	2	3	4	5

PART 2 – STUDENTS WHO DO NOT PERFORM SO WELL
STUDENT 1

What does this student do?

On a 5-point scale, how do you rate this students' reading and analysis of the material? Justify your answer.

Very poor	Poor	Average	Good	Very good
1	2	3	4	5

STUDENT 2

What does this student do?

On a 5-point scale, how do you rate this students' reading and analysis of the material? Justify your answer.

Very poor	Poor	Average	Good	Very good
1	2	3	4	5

STUDENT 3

What does this student do?

On a 5-point scale, how do you rate this students' reading and analysis of the material? Justify your answer.

Very poor	Poor	Average	Good	Very good
1	2	3	4	5

Appendix B

Description of the EMMEs (rubric)

EMME	Strategy modeled	Description
Students who perform well		
1	SERP inspection	A student types keywords on Google, inspects a SERP page from top to bottom to decide which website to access, reading at a normal pace all the page titles and some information from the snippets. The student ends up clicking on a relevant page at the bottom of SERP after a review of the SERP titles
2	Identification of source information	A student clicks on a SERP result, enters the webpage, looks at the webpage logo (a mineral water brand), reads the text once at normal pace, and reads the author information provided below the text (commercial director)
3	Deep reading of trustworthy and relevant pages	A student clicks on a SERP result, enters the webpage, looks at the webpage logo (institutional page – consumers’ association), reads the text twice at normal pace, and reads the author information provided below the text (doctor). They go back and continue to inspect the SERP.
4	Skimming of less trustworthy pages	A student clicks on the second SERP result, enters the webpage, inspects author information on the left of the text (picture, name and occupation), then looks at the webpage logo (popular forum), quickly skims the text and reinspects author information before leaving the page.
5	Quickly abandoning topically unrelated pages	A student clicks on a SERP result, enters the webpage, looks at the webpage logo (waterpark i.e., commercial service unrelated to the task), the image, reads the text title, and abandons the page without reading the text.
Students who do not perform so well		
6	SERP inspection	A student starts a search using the strings “fuels and climate change”. The SERP results appear. Student clicks on the first website without previously inspecting the other search results
7	Identification of source information	A student starts reading the webpage from the title of the text without looking at the banner, logo or author information
8	Deep reading of trustworthy and relevant pages	The student looks at the webpage banner but does not fully read the text even though it is reliable

Appendix C

Patterns of misinterpretation

Strategy 2 – Identification of source information		
Participant	Example	
1	EMME01	I don't think it was a good idea to search information from a website company because it's going to tell you what you want to hear to buy its product, even if it's not necessarily true.
2	EMME03	I don't think he chose a useful article because it's written in a brand's web "Fontvella". Of course this brand will write good information about the product that is being sold.
3	EMME06	The strategy was poor because they didn't realize images were worthless, some sentences were ignored and there was no intention of making an overview of the content to find the most important input.
4	EMME08	The student only skimmed through the text and paid more attention to the name and the logo. Most importantly, his chosen resources was of questionable trustworthiness.
5	EMME10	The reason why I selected this answer is because they don't seem to read one sentence and spend too much time looking at the webpage name and the author of the text.
6	EMME16	The student clicks on a website that will probably lie because it is written by commerce so they have the intention to sell.
7	EMME20	POOR He opens a new review and reads it slowly, understanding what he is reading. He needs to compare to other places because Iberdrola wants to say positive things to have new interested customers.
8	EMME30	The title seems to be the relevant information for him. He don't pay attention to the important information: the text.
9	EMME52	The student spent a lot of time reading the title of the website and when it was time to read the main point, he didn't read all of it.
Strategy 4 – Skimming of less trustworthy pages		
Participant	Example	
1	EMME01	Looking for information on this kind of website is not reliable because everyone can post whatever they want even if it's not true.
2	EMME02	He is not focused on reading the information, he is looking for other things.
3	EMME03	I don't think he did a good search because the source of the article is not very reliable, it is just based on someone's experience.
4	EMME24	The student spent way more time checking the information about the content's author than on actually reading the material.
5	EMME30	The text was read very quickly. The details were too observed.
6	EMME32	The information that he was reading was an opinion without specific knowledge.
7	EMME34	He/she chooses a website where all the people can write and he gives more attention to the profile than to the text.
8	EMME36	The student is searching for personal experiences.

9	EMME54	It is very poor because you cannot trust a forum where random people give their opinions with no scientific evidence.
---	--------	---

Strategy 5 – Quickly abandoning topically unrelated pages

Participant	Example
1 EMME01	Even if he/she hasn't read the text I think it was a mistake to enter the website on the first place.
2 EMME02	He is looking for the wrong information and he doesn't read it completely.
3 EMME03	His research was very bad because the article is not even related to the information he's looking for.
4 EMME14	The student chooses a source which was not related with the question, maybe they should have read the headlines more carefully.
5 EMME20	He opens a link that goes to a theme park, that has nothing to do with his search.
6 EMME32	He doesn't get the information that he wants because he is overiewing a commercial text.
7 EMME44	I think the student filtered unnecessary information. By not reading the text properly, they might have missed important information.
8 EMME46	Too fast. He/she does not look for any detail information. Even though we assume he/she does it pretty good because the website doesn't have any interest information for the topic.
9 EMME48	It's not good because the student reads the paper but the main topic is not correct to answer the question's task.

Strategies 4 and 5 – skimming misinterpreted

Strategy 4

1	EMME30	The text was read very quickly. The details were too observed.
2	EMME42	This student only reads the second article and very fast.
3	EMME52	It is true that he checked the source of information but he read the text extremely fast so I don't think he fully understood.
4	EMME56	I think the reading and analysis of the material is poor because it's impossible for a person to get information from a text when reading at that speed.
5	EMME58	It's ok that you learn about who writes the article but it's worse to read it quickly.

Strategy 5

6	EMME02	He is looking for the wrong information and he doesn't read it completely.
7	EMME16	He/she doES not finish reading and attention goes to the picture.
8	EMME24	The student didn't even read the content of the page and was acting chaotically.
9	EMME26	The student only reads the title.
10	EMME44	I think the student filtered unnecessary information. By not reading the text properly, they might have missed important information, though.
11	EMME46	Too fast. He/she does not look for any detail information. Even though we assume he/she does it pretty good because the website doesn't have any interest informatin for the topic.

12	EMME52	The student read the title and the first line of the text and saw pretty fast that this wasn't what he was looking for but I think he should have read a little more before exit the website.
13	EMME54	It is very important to read at least the titles of the links and decide if they are in the topic of the research.
14	EMME56	This analysis was very poor because the student didn't even read the text.
15	EMME58	You must read more information if you want to do a good research.

This preprint was submitted under the following conditions:

- The authors declare that the necessary Terms of Free and Informed Consent of participants or patients in the research were obtained and are described in the manuscript, when applicable.
- The authors declare that the preparation of the manuscript followed the ethical norms of scientific communication.
- The authors declare that they are aware that they are solely responsible for the content of the preprint and that the deposit in SciELO Preprints does not mean any commitment on the part of SciELO, except its preservation and dissemination.
- The authors declare that the data, applications, and other content underlying the manuscript are referenced.
- The deposited manuscript is in PDF format.
- The authors declare that the research that originated the manuscript followed good ethical practices and that the necessary approvals from research ethics committees, when applicable, are described in the manuscript.
- The authors declare that once a manuscript is posted on the SciELO Preprints server, it can only be taken down on request to the SciELO Preprints server Editorial Secretariat, who will post a retraction notice in its place.
- The authors agree that the approved manuscript will be made available under a [Creative Commons CC-BY](#) license.
- The submitting author declares that the contributions of all authors and conflict of interest statement are included explicitly and in specific sections of the manuscript.
- The authors declare that the manuscript was not deposited and/or previously made available on another preprint server or published by a journal.
- If the manuscript is being reviewed or being prepared for publishing but not yet published by a journal, the authors declare that they have received authorization from the journal to make this deposit.
- The submitting author declares that all authors of the manuscript agree with the submission to SciELO Preprints.