



Delft University of Technology

Effects of Experts' Annotations on Fashion Designers Apprentices' Gaze Patterns and Verbalisations

Coppi, Alessia Eletta; Oertel, Catharine; Cattaneo, Alberto

DOI

[10.1007/s12186-021-09270-8](https://doi.org/10.1007/s12186-021-09270-8)

Publication date

2021

Document Version

Final published version

Published in

Vocations and Learning: Studies in vocational and professional education

Citation (APA)

Coppi, A. E., Oertel, C., & Cattaneo, A. (2021). Effects of Experts' Annotations on Fashion Designers Apprentices' Gaze Patterns and Verbalisations. *Vocations and Learning: Studies in vocational and professional education*, 14(3), 511–531 . <https://doi.org/10.1007/s12186-021-09270-8>

Important note

To cite this publication, please use the final published version (if applicable).
Please check the document version above.

Copyright

Other than for strictly personal use, it is not permitted to download, forward or distribute the text or part of it, without the consent of the author(s) and/or copyright holder(s), unless the work is under an open content license such as Creative Commons.

Takedown policy

Please contact us and provide details if you believe this document breaches copyrights.
We will remove access to the work immediately and investigate your claim.



Effects of Experts' Annotations on Fashion Designers Apprentices' Gaze Patterns and Verbalisations

Alessia Eletta Coppi¹ · Catharine Oertel² · Alberto Cattaneo¹

Received: 5 March 2020 / Accepted: 15 April 2021 / Published online: 15 July 2021
© The Author(s) 2021

Abstract

Visual expertise is a fundamental proficiency in many vocations and many questions have risen on the topic, with studies looking at experts and novices differences' in observation (e.g., radiologists) or at ways to help novices achieve visual expertise (e.g., through annotations). However, most of these studies focus on white-collar professions and overlook vocational ones. For example, observing is uttermost important for fashion designers who spend most of their professional time on visual tasks related to creating patterns and garments or performing alterations. Therefore, this study focuses on trying to convey a professional way to look at images by exposing apprentices to images annotated (e.g., circles) by experts and identifying if their gaze (e.g., fixation durations and gaze coverage) and verbalisations (i.e., images descriptions) are affected. The study was conducted with 38 apprentices that were exposed to sequential sets of images depicting shirts, first non-annotated (pre-test), then annotated for the experimental group and non-annotated for the control group (training 1 and training 2), and finally non-annotated (post-test). Also, in the pre and post-test and in training 2 apprentices had to verbally describe each image. Gaze was recorded with the Tobii X2–60 tracker. Results for fixation durations showed that the experimental group looked longer in the annotated part of the shirt in training 1 and in the shirt's central part at post-test. However, the experimental group did not cover a significantly larger area of the shirt compared to control and verbalisations show no difference between the groups at post-test.

Keywords Visual expertise · Gaze · Annotations · Fashion designers

✉ Alessia Eletta Coppi
alessia.coppi@iuffp.swiss

Catharine Oertel
C.R.M.M.Oertel@tudelft.nl

Alberto Cattaneo
alberto.cattaneo@iuffp.swiss

¹ Swiss Federal Institute for Vocational Education and Training, Via Besso 84–86, 6900 Lugano, Switzerland

² Delft University of Technology, Interactive Intelligence Group, Van Mourik Broekmanweg, 6, 2628 Delft, Netherlands

Introduction

The professional world is in constant evolution and, inevitably, the field of education must remain up-to-date by attempting to prepare students for a hyper-specialised and technology-driven world. Numerous transversal skills are necessary to thrive in an ever-changing professional environment (Pang et al., 2019)—both personal and interpersonal or emotional skills (Srivastava, 2013), as well as other more specific skills like visual expertise, are indispensable for certain professions (Jasani & Saks, 2013; Kok et al., 2017); Naweed & Balakrishnan, 2014; O'Halloran & Deale, 2006).

Observation has long been studied from various perspectives. From a sociological point of view, it can be considered interconnected with expertise: observation is a shared perception of the world that is visible in specific shared actions (e.g., using coding schemes) within a professional community (Goodwin, 1994). From a more cognitive perspective, it can be viewed as a set of visual skills that are specific to each task and domain and require interaction between perceptual and cognitive skills (Raveslout et al., 2012).

Examples of such skills include visual search and visual information processing to detect and interpret lesions in X-rays (see Raveslout et al., 2012) or to recognise specific patterns, such as presence of symmetry, colours, and shapes in radiology (Shapiro et al., 2006). The ability to use global and local processing of image features in arts (Chamberlain & Wagemans, 2015) or the skills to identify specific students' behaviours in class for teachers (Stürmer et al., 2013) constitute other examples of such skills. The importance of visual expertise makes it a topic that is studied in numerous contexts, such as medicine (Naghshineh et al., 2008), architecture (Styhre & Gluch, 2009), and sports (Kredel et al., 2017). Most of these studies focus on understanding the differences between experts and novices, either to distinguish where they focus their attention while looking at professional material or investigate how to improve observation skills of novices. In this latter case, the strategy is often to redirect novices' attention using annotations (or signals or cues) in the form of arrows, text, audio, colours, etc. (Schneider et al., 2018).

However, studies on observation have rarely been explored within vocational education, although some examples are worth investigating. For instance, fashion designers are certainly among those vocational professions that strongly rely on visual expertise. In Switzerland, fashion designers have a dedicated full-time and dual (part-time school and part-time in the workplace) Vocational Education and Training (VET) course. During their three years' training, they have to engage with tasks like production of sketches, production and alteration of patterns, knowledge and use of fabrics, styling of clothing, and production and alteration of garments both by hand and using professional machines. Apprentices have to acquire a wide range of skills to work in the sartorial world, of which visual expertise is one of the most important since they have to learn how to observe different representations of garments such as photos, technical sketches, or patterns (Coppi et al., 2019) and garments themselves. More specifically, looking at their

training plan (State Secretariat for Education, Research and Innovation, Fashion Designers' Training Regulation, 2018), it emerges that in order to become an acknowledged professional they must develop the competence to observe different details (e.g., specific manufacturing, types of pockets, garment opening types, types and quality of fabric), defects (e.g., defects of proportions, wearability, manufacturing quality) and characteristics of the customer's body (e.g., posture, shape irregularities to hide or enhance) (see Caruso et al., 2017). At the same time, looking at their learning activities, they have to develop their competence to make their observation skills explicit when professionally describing a garment to become an acknowledged professional (Caruso et al., 2019).

This paper aims to address the above-mentioned research gap by investigating how visual expertise can be fostered in apprentice fashion designers. This profession can be ideal for creating experimental set-ups specific to the many tasks that apprentices have to perform and to convey a professional way to look at profession-specific images using cueing methods. In particular, this study intends to tackle the following two research questions:

- Can exposure to annotations convey a professional way to look at images in terms of fixation durations and gaze coverage?
- Can exposure to annotations affect the way apprentice describe professional images?

Visual Expertise and Eye Movements

In order to understand how to convey a professional way of looking at images, some background information is needed about eye movements' parameters and how they relate to expertise.

Among the main types of eye movements, the most used are fixations and saccades. Fixations are the period of time when the eye is relatively still (Holmqvist et al., 2011) and reveal where a participant's attention is allocated on a specific stimuli (or parts of the stimulus). Saccades are quick and simultaneous movements of the eyes between fixations in the same direction (Cassin et al., 1984) and they help in understanding the direction of the gaze and if targets attract attention. These types of eye movements are fundamental to research on observation since difference in expertise can be seen in both of these parameters but mainly in fixation durations.

For instance, expert gymnastic coaches show fewer but longer fixations on an athlete's body compared to novice coaches (Moreno et al., 2002); similarly, expert swimming coaches spend more time fixating on swimmers' movements (e.g., body-roll) compared to their novice colleagues and specific body parts (e.g., hands and head) when the swimmer is entering the water (Moreno et al., 2006). Similarly, climbing experts had fewer but longer fixations on a climber's body (core and feet and feet placements) with more and shorter fixations on other parts of the body (hands) compared to novices (Mitchell et al., 2020).

In the medical field, experts spend their time fixating on areas of high saliency compared to novices, while students have generally prolonged saccades on the entire image (Krupinski et al., 2006). More recently, Krupinski et al. (2014) trained

dermatologists in reading benign and malign close-up dermoscopy images and showed that a dermatologist tends to have more efficient search with fewer fixations and longer dwells ('one visit in an area of interest, from entry to exit'; (Holmqvist et al., 2011, p. 190). Further, in radiology, Wood et al. (2013) found that experts are quicker than novices in identifying fractures and spent more time fixating on them. In addition, experts draw conclusions based on observed abnormalities instead of focusing directly on pathologies like novices do (Jaarsma et al., 2014).

In biology, Jarodzka et al. (2010) indicated that in a fish locomotion task, novices had longer viewing times on irrelevant areas, while experts had longer gaze durations on the features of fish and also provided richer verbal descriptions of the fishes' movements. In education, expert teachers were better able to perceive class interactions by generally visiting and revisiting specific areas, gazing at a student's body, and at sections of the class with more physical and verbal interaction between students (Wolff et al., 2016).

Generally, regardless of the field, it seems that visual expertise appears to rely on three main elements: 1) being able to identify more quickly the task-relevant areas of the image; 2) spending more time observing relevant features (e.g., fracture in X-ray) and less time on irrelevant ones (e.g., tender tissues in X-ray); and 3) being more cognitively activated, that is, showing higher content evaluation and understanding of the images viewed. Also, as another measure for expertise, 4) the expert should be able to verbalise more technical terms and more features related to the task compared to novices.

Strategies for Improving Visual Expertise and the Role of Annotations

Research did not stop at identifying the differences between novices and experts but also attempted to understand how to improve observation in novices to make them observe more like an expert. Most of these studies draw their theoretical basis from the concept of signalling that, according to Mayer (2002), is the process of presenting cues in an effective manner so that students can easily select—and then process—instructional material better. Signals are placed in the most significant areas of the learning material (e.g., text or image) to redirect students' attention and to make certain areas more salient than others. This can result in lower cognitive load, more focused attention, and better comprehension and retention (Dodd & Antonenko, 2012; Richter et al., 2016).

Cues, signals, or annotations can have multiple forms—colours, text-picture reference, gestures, labels, flashlights, graphic organisers, and diagrammatic elements such as lines, blobs, boxes, crosses, arrows, circles, and transparent overlay (Heiser & Tversky, 2006; Shin & Park, 2019; van Gog, 2014). For example, coloured indicators can help participants with low spatial skill in solving mental rotation tasks (Roach et al., 2019). Also, translucent colour patches were used in golf training videos to show novices the correct movement to adopt to improve their performance (D'Innocenzo et al., 2016). Others used a spotlight method and applied shades of colour to specific sections of a medical image and were able to induce in novices longer fixations on cued areas compared to non-cued ones (De Koning et al., 2007).

Moreover, the use of lines or arrows have proven useful. For example, arrows applied to an animation displaying the functioning of a piano increased the attention towards relevant sections and resulted in better comprehension of the movements of the piano keys (Boucheix & Lowe, 2010). In others, circles and colours were used to improve the diagnostic performance and visual processing of novice radiologists by showing participants radiographies with relevant parts highlighted with yellow circles, shading on irrelevant areas, and use of colours to highlight the lines of the fractures (Scheiter et al., 2019).

Assuming that experts (e.g., radiologists) tend to observe the most important areas (e.g., lesion) of an image (e.g., X-ray), other scholars used dots positioned on the material to replicate the eye movements of an expert as modelling cues. In this manner, learners are exposed to an expert's gaze pattern under the hypothesis that they will implicitly learn where to look and, hopefully, replicate the pattern in other materials. A pioneering study in this area was performed by Grant and Spivey (2003), and exposed participants to the annotations modelled on the gaze of an expert to highlight sections of a problem-solving task, helping students find the solution. Other scholars asked novice marine biologists to watch and listen to a video showing a fish moving in the water and identify the types of movement it made. The video contained annotations based on the eye movements of an expert that helped novices focus on a section of the body of the fish to identify the movements of the fish (see Jarodzka et al., 2010, 2013). Further, Sharma et al. (2015) superimposed teachers' gazes on an educational video and helped students have a more linear experience, better understanding of the teacher's gestures, and higher engagement. Similarly, Mason, Pluchino, Tornatora, and Ariasi (2013) showed participants the models' eyes in the form of red dots on a text with graphics and noticed that learners spent more time re-fixating on the relevant parts of the visualisation while reading and performed better than control in a recall test.

Hypotheses

Based on the above-mentioned literature, annotations appear to be an effective technique for promoting visual expertise, often combined with eye tracking techniques. Consequently, this study was developed under the assumption that apprentices exposed to annotated pictures will learn to explore images more thoroughly and pay greater attention to the annotated areas compared to control participants exposed to the same images without any annotation, thereby producing a richer verbal description of the same pictures.

Therefore, we hypothesize that:

- H1) the fixation durations will be longer (more attention allocated) in specific areas of the shirt but shorter in irrelevant and non-annotated parts of the image in the group exposed to annotations (experimental) compared to a group not profiting from annotation (control) during the treatment;
- H2) the fixation durations in relevant areas will be longer after the treatment compared to the control group in the different areas of the shirt;

- H3) the gaze coverage of the image will be larger in the experimental than in the control group;
- H4) the experimental group will be able to produce a significantly higher number of details in the verbalized description of the images after the intervention.

Methods

Sample

The sample comprises 38 apprentices randomly assigned into two groups, 20 for the experimental ($Mage=16.8$, $SDage=1.21$, all females) and 18 for the control group ($Mage=17.56$, $SDage=3.83$, all females). The apprentices were enrolled at the beginning of the second year of their three-year certificate program and were recruited from two vocational schools in Switzerland, which provides a well-established fashion design curriculum.

Original Task and Design of the Experiment

Activities that rely on visual expertise are part of the daily practice of apprentices and teachers provide students with opportunities to train this skill. For this reason, we took inspiration for the design of our experiment from a didactical scenario that students perform weekly in their course since the first year in which focus on a wide range of skirts and trousers, then on shirts in the second year and coats and more complex garments in the third. This task is interesting 1) for the manner in which teachers use annotations to direct a student's attention and 2) for the request to analyse the garment. In this exercise, students have to observe the image 'Fig. 1' (left) and provide a full description comprising possible measurements and number and type of parts that compose a garment. Thereafter, the teachers provide the same image with annotations such as lines and arrows or colours and text and initiate a conversation with the students that reviews the entire image 'Fig. 1' (right).

Stimuli and Annotations

Stimuli were selected with the help of the teachers who identified 20 photographs depicting shirts with a range of peculiarities related to shapes, style, details, and parts. These images were then distributed in four subsets of five images each, to be used for the pre-test (5 images), training 1 (5 images), training 2 (5 images), and post-test (5 images). The distribution was made by the teachers in order to make the four subsets equal with respect to their general complexity. Teachers also provided detailed descriptions of each image, highlighting the relevant areas of the garment such as sleeves, stitching, pockets, neck, cuffs, hem, buttons, and the opening (an open section at the top of a garment for the neck of the wearer). These are areas of high relevance and are the ones on which annotations were positioned in the training sessions of the experiment. Annotations took the form of circles, squares, and

Pattern making course

TEST N. 3

Name and surname: _____ Time available: 90 min.

Alter the skirt using the following information and produce a pattern with a belt.

CLIENT MEASURES:

Hips 37 cm
Waist 47.5 cm

CALCULATIONS:

Back



74.5 cm

48 cm

DESCRIPTION:

GOOD LUCK!

Pattern making course

CORRECTION - TEST N. 3

Alter the skirt using the following information and produce a pattern with a belt.

CLIENT MEASURES:

Hips 37 cm + 1 cm = 38 cm
Waist 47.5 cm + 1 cm = 48.5 cm

CALCULATIONS:

Width: 74.5 cm - Hips 48.5 cm = 26 cm
Calculation for division: (distance between the lines) 26 cm : 5 (division) = 5.2 cm
Open hips 5.2 cm : 2 = 2.6 cm
Position of opening lines: hips 48.5 cm : 6 (spaces) = 8.08 cm - 8.1 cm



Back

Width

74.5 cm

48 cm

Zip centre right and left

Stitches on the hips

Middle section

NO darts

Straight belt

DESCRIPTION:

Skirt with wide/flare line

Zip on the inside of the left hip

Middle length

V line with a belt

Yielder: silk satin

1. Observe the image → Identify the type of skirt

2. Observe the details (e.g., openings, etc...)

3. Add ease

4. Recognise the measure for right and left length/width

5. Execution of calculations

6. Describe the image

Fig. 1 Examples of the activity before completion (left) and corrected by the teacher (right)

rectangles and these specific shapes were selected to redirect students' attention to defined and specific areas and keep their gaze within those delimited sections while distinguishing and highlighting features present in adjacent areas ('Fig. 2'). Moreover, the position of the annotations was determined not only to highlight

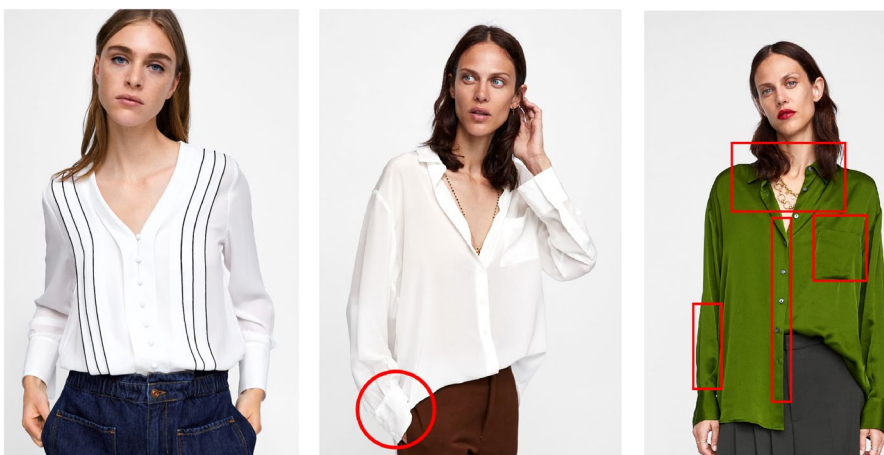


Fig. 2 Annotations placement: non-annotated (pre-test and post-test, left) and annotated image (training 1 and 2 for the experimental group, centre and right)

the corresponding areas but to foster a systematic and comprehensive way to observe the images

Apparatus

The Tobii Pro X2-60 Eye tracker was used to collect the students' eye movement. The tracker is able to sample data binocularly at 60 Hz. The tracker holds a gaze accuracy of 0.4° (binocular and monocular), gaze precision of 0.34° (monocular) and 0.45° (binocular). Gaze measures were classified with the default setting of Tobii Pro Lab (velocity-based algorithm: peak velocity $30^\circ/\text{s}$, min. fixation duration 60 ms). Stimuli were presented via the Tobii Pro software (Version 3.4.8) connected to an HP laptop and a 15-inch PC screen (HP ZR2440w) with the maximum resolution of 1920×1200 pixels. The eye tracker, screen, and laptop were placed on an ergonomic table that could be adjusted by height. Also, a chair with resting arms was placed at the recommended distance (65 cm) in front of the table and both the seat and table were placed at an appropriate height to capture the gaze of each participant. A head stabilisation system was not required.

Procedure

Before the beginning of the sessions, the researchers set up all the tools and adjusted the room: tables and chairs were positioned, the brightness from the windows was adjusted using blinders, and the positions of the chair, table, and computer screen were marked with tape. Afterwards, participants were welcomed in the room only one at a time and were asked to find the most comfortable position on the chair. Basic information about the experiment and the functioning of the eye tracker was provided, and the participants were instructed to avoid movements such as tilting or turning the head, looking away from the screen or at the researcher, and obstructing the eye tracker in any way. Then, the experimenter initialised the 9-point calibration procedure. Due to time constraints imposed by the school and the length of the

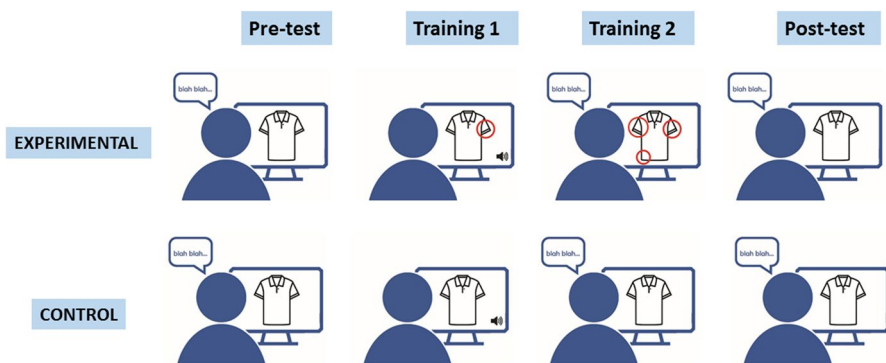


Fig. 3 Schematized presentation of the procedure with the two groups

design, calibration was repeated only in case of loss of connection between the PC and the eye tracker.

The specific procedure of the experiment was performed in the following manner (see Fig. 3): in the pre-test, each participant (experimental and control) was exposed to a set of five images depicting garments displayed on a PC. Each picture was displayed for 40 s (plus 20 s of a black screen before the next image), during which the participants were asked to verbally describe (concurrent think-aloud) the garment.

In training 1, participants of the experimental group were asked to look at a series of five images for a total of 03.27 min; each of them was completed with visual annotations and an audio description inserted by the teacher. The visual annotations entered progressively alongside the related piece of the audio description and then faded out in a manner in which only one annotation was present on the screen at the same time. Participants in the control group were asked to look at the same set of images, but no visual annotations were provided to them.

In training 2, the experimental group was presented with another set of five images and asked to observe each of them and to describe it verbally (each image was displayed for 40 s each, followed by 20 s of a black screen before the next image). Each image was annotated, but in contrast to training 1, the annotations were present on the image for the entire duration of 40 s. In contrast, the control group was exposed to the same set of images, but without any annotations, and was similarly asked to describe them verbally.

In the post-test, all the participants (experimental and control) were asked to observe a new set of 5 images for 40 s each (plus 20 s of a black screen before the next image) and to provide a verbal description as they had done in the *pre-test* phase. No annotation was present in this case. Each image was initialised manually by the experimenter after warning the participant that a new image was going to be displayed.

The whole procedure lasted about 30 min: although one could assume that these kinds of intervention might need more time to be effective, we had to limit the duration of the intervention to make it feasible for the apprentices and the school.

Measures

Fixation durations: this measure comprises the duration of fixations on each of the images presented in pre-test, training 1, training 2 and post-test.

Gaze coverage: a metric used in the area of eye tracking but defined in different ways (see Holmqvist et al., 2011; Van der Gijp et al., 2017). Coverage has been used in a limited number of studies that pointed out that experts and novices differ in the amount to which they observe images or videos (e.g., Jaarsma et al., 2014). In our case it was defined as “*the proportion of the image which is gazed at by the individual at least once. The measure was computed by retrieving the heatmaps for each given image, estimated over the duration of the stimuli. Then, we thresholded the heatmaps based on a fixed value and computed the gaze coverage as the ratio of pixels which were above the threshold over the total amount of pixels in the image.*”

In our experiments, the threshold was set to 0.005, given that the heatmap values were normalised so that the maximum value was 1” (see Oertel et al., 2019, p. 390).

Verbalisation score (concurrent think-aloud): to calculate the score, audio of the participants’ verbalisations was extracted and manually transcribed. The score was based on the written professional descriptions of all the images provided by the teachers involved in co-designing the study. A list of 52 features (e.g., opening, neck, bottom, cuff, hem, pockets, ribbon, fit, ruffles, stitching, seam, etc.) were derived and these served as our coding scheme. Participants’ scores were calculated by giving 1 to each detail correctly mentioned (i.e., present in the participant’s description and in the teacher’s description) and 0 to any missing or incorrectly mentioned detail. Since the number of details included was not the same for each image (they ranged from 8 to 17), a weighted score was calculated for each image and then an average weighted score was calculated per group for the different phases of the test.

Areas of Interest (AOIs) for Fixation Duration Analyses

To be able to work with fixation durations it is necessary to connect the eye tracking data with the visual stimuli by using AOIs that, in our case, were manually drawn directly in Tobii Pro Studio. The AOIs have a semantic composition defined by the



Fig. 4 Example of pre-test and post-test AOIs (left) and with the AOIs (right): ‘Sleeve’ (polka dot), ‘Central’ (black), ‘Neck’ (grey and white stripes) and ‘Details’ (black and white stripes)

experts (see Holmqvist et al., 2011) and in pre and post-test (Fig. 4) AOIs were placed on the section of the shirt (central, neck, sleeve and details) deemed most relevant for the profession in a specific garment since they are observed to understand how the garment is constructed and then how to produce a correct pattern needed for assembling the final garment. These areas are:

- the ‘central’ part of the shirt is observed to identify the position and type of opening, which characterizes most a shirt; it can also contain cuts, ruffles or folds, pockets and more. Also, this section can dictate the desired fit of the shirt;
- the ‘neck’ is looked at to understand the type and the connection to the opening of the shirt (for instance V openings connect directly to the neck while openings on the side with a zip often require a completely different pattern);
- the ‘sleeves’ are important to observe to understand where the central section can be connected (sewed) to the shoulder, to identify the type of sleeve, shoulder features, the type and closure of the cuff (as well as related details) and the desired fit;
- the ‘details’ include embellishments positioned on different areas of the shirt such as ribbons, embroideries, piping (passepoil)—but also stitching or buttons that can be used for merely aesthetic reasons. Details can influence the pattern and their inclusion might require changes in its design.



Fig. 5 Example of OAIs placed on the shirt in training 1 and training 2—The AOIs ‘Sleeve’ (polka dot), ‘Neck’ (grey and white stripes), ‘Central’ (black), ‘Details’ (black and white stripes) and ‘Shirt’ (grey)

Conversely, in training 1 and 2, the AOIs (Fig. 5) were placed on the annotations and the non-annotated part of the shirt. Training 1 and training 2 contains the same AOIs of pre and post-test (central, neck, sleeve) with inclusion of the AOI 'Shirt' placed on top of the areas of the shirt where no annotations had been put by the teachers and the absence of the AOI 'Details' for training 1.

These AOIs are also important for contextualising the participants' verbalisations: indeed, they allow to put in relation what the participant has "looked at" and what has actually "seen" in those fixations. The combination of the AOIs data and the verbalisations allows for a more complex understanding of the apprentices' observational experience, helping to connect the perceptual action of observing a specific section with the processes of the working memory allowing understanding of the reasoning and knowledge behind both novices and experts (Jarodzka & Boshuizen, 2017).

Missing Data and Exclusion Criteria

Due to technical issues, gaze data was lost for two participants in the experimental and one in the control group while audio–video was lost for one participant in the control group. Therefore, the analyses on gaze coverage and fixation durations were performed on 35 participants and the verbalisations on 37. No participant was excluded due to eye tracking calibration issues since calibration was repeated until a good signal was reached.

Data Preparation and Analysis for Fixation Durations and Gaze Coverage

The software IBM SPSS 26 was used to carry out all analyses for fixation durations and gaze coverage. General linear mixed effects models (GLMMs) were used to assess differences between the groups in both fixation duration per AOIs and for gaze coverage. GLMMs were used because they can be an effective method of analysis in case of repeated measures analysis of variance or in many other cases such as with unbalanced designs (i.e., with unequal number of participants within each level of a grouping variable), incomplete data (e.g., with missing observations at one or more time points for numerous participants), non-independence among observations (i.e., when the same task is repeatedly administered to each participant; (see Baayen et al., 2008; Judd et al., 2012). The GLMMs for fixation durations and gaze coverage were performed with a design with crossed random effects with nested observations within participants. Fixation durations data was first extracted from Tobi Pro Studio and imported in Excel for cleanup, and then it was transformed from wide to long format in SPSS. The GLMMs for fixation durations were run with condition and test as fixed factors and subjects and AOIs as random factors. Also, a Bonferroni correction was applied to the data. Gaze coverage database was obtained with a script and was also transformed in long format in SPSS where data were analysed with condition and test as fixed factor and subjects and gaze coverage as random factors. A Bonferroni correction was applied in this case too.

Results

Fixation Durations

A series of GLMMs were conducted to identify the effects of presence of annotations on fixation durations on the AOIs for training 1 and training 2.

In the AOI 'Central', results showed a significant main effect of condition [$F(1, 33)=7.649, p<0.01$] and a significant interaction condition*test [$F(1, 305)=6.008, p<0.025$]. Planned post-hoc comparisons of condition indicated a significant difference between the groups [$F(1, 33.00)=7.649, p<0.01$] while comparisons of the interaction showed a significant difference between the groups in training 1 [$F(1, 52.508)=12.829, p<0.001, d=0.584$].

In the AOI 'Sleeve', results showed a significant interaction condition*test [$F(1, 305.000)=79.621, p<0.001$]. Planned post-hoc comparisons of the interaction showed a significant difference between the groups in training 1 [$F(1, 89.901)=41.016, p<0.001, d=1.55$] and in training 2 [$F(1, 89.901)=25.128, p<0.001, d=0.589$].

In the AOI 'Neck', results indicated a significant interaction condition*test [$F(1, 305)=3.604, p=0.05$]. Planned post-hoc comparisons of the interaction indicated a significant difference between the groups in training 1 [$F(1, 59.610)=4.197, p<0.05, d=0.329$].

In the AOI 'Shirt' (non-annotated part of the shirt), results showed a significant main effect of test [$F(1, 8.000)=6.349, p<0.05$] and a significant interaction condition*test [$F(1, 305)=6.196, p<0.025$]. Planned post-hoc comparisons of test indicated a significant difference between the trainings [$F(1, 8.001)=6.349, p<0.05$] while comparisons of the interaction showed a significant difference between the groups in training 2 [$F(1, 46.884)=3.979, p=0.05, d=0.478$].

No difference was identified in the AOI 'Details' in both training 1 and training 2.

Further, other GLMMs were conducted to identify the effects of presence of annotations on fixation durations on the AOIs for the pre and post-test.

In the AOI 'Central', results showed a significant interaction condition*test [$F(1, 305)=6.321, p<0.025$]. Planned post-hoc comparisons of the interaction indicated a significant difference between the groups at post-test [$F(1, 53.497)=5.231, p<0.05, d=0.579$]. However, results in the others AOIs 'Neck', 'Sleeve' and 'Details' showed no difference between the experimental and control group.

All means and standard deviations are available in the table below (Table 1).

Gaze Coverage

A GLMM was conducted to compare the effect of presence of annotations on gaze coverage. Results showed that the interaction between condition and gaze coverage is not significant [$F(1, 236.091)=3.747, p=0.05$]. Planned post-hoc comparisons indicated a significant difference [$F(1, 48.289)=4.510, p<0.05, d=0.451$] between the experimental and control groups at the post-test. The mean values indicate that

Table 1 Means and standard deviations (SDs) of the groups on all tests

	Pre-test		Training 1		Training 2		Post-test	
	<i>M(SD)</i>	<i>M(SD)</i>	<i>M(SD)</i>	<i>M(SD)</i>	<i>M(SD)</i>	<i>M(SD)</i>	<i>M(SD)</i>	<i>M(SD)</i>
	EXP	CON	EXP	CON	EXP	CON	EXP	CON
Central	3.90 (1.56)	3.92 (1.58)	2.62 (1.53)	1.84 (1.12)	1.55 (1.06)	1.26 (0.82)	3.72 (1.62)	3.13 (1.26)
Neck	5.77 (3.71)	6.12 (4.00)	2.14 (1.70)	1.63 (1.37)	3.11 (1.53)	3.08 (1.81)	3.54 (2.31)	3.77 (2.61)
Sleeve	1.20 (0.630)	1.19 (0.694)	1.75 (1.09)	0.723 (0.610)	1.09 (1.21)	1.90 (1.53)	1.34 (0.734)	1.28 (0.726)
Details	1.43 (1.71)	1.46 (1.75)	-	-	2.38 (1.77)	2.19 (1.79)	2.03 (2.43)	1.81 (1.93)
Shirt	-	-	18.14 (7.13)	18.13 (6.98)	11.56 (4.95)	13.98 (5.17)	-	-

in post-test, the experimental covered a larger portion of the image compared to control (respectively $M_{exp} = 10.95$, $SD = 0.287$ vs $M_{con} = 10.82$, $SD = 0.289$).

Verbalisations

To identify a possible difference between the groups, a repeated-measure ANOVA was conducted to compare the effects of condition on the verbalisation score (number of enunciated details) in both the pre-test and post-test. The results revealed a significant main effect of verbalisation score [$F(1, 35) = 7.681$, $p < 0.01$, $\eta_G^2 = 0.769$]. However, the between-subject effect was not significant, $F(1, 35) = 3.423$, $p > 0.05$, $\eta_G^2 = 0.436$] as well as the interaction [$F(1, 35) = 0.640$, $p > 0.05$, $\eta_G^2 = 0.122$] (Fig. 6).

Discussion

Are There Significant Differences in Terms of Fixation Duration Between the Groups?

Our main hypothesis was that the experimental group would look longer at the annotated AOIs ('Neck', 'Central', 'Sleeve', and 'Details') and less at non-annotated areas ('Shirt') compared to the control group. Also, that at post-test the experimental group would look longer at the relevant areas of the shirt compared to the control group. Results of training 1 indicated that the experimental group looked longer at the AOIs

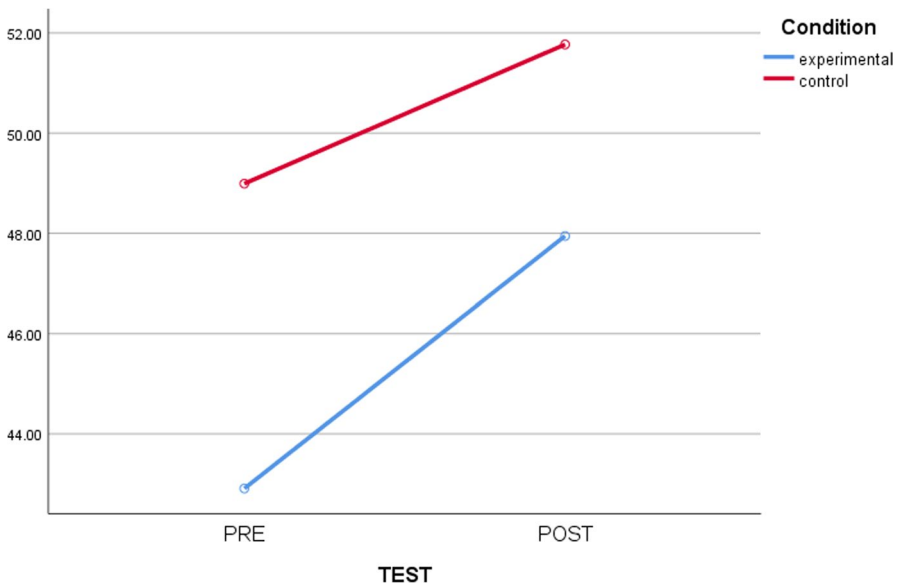


Fig. 6 Number of details mentioned by the experimental and control groups in the pre-test and post-test

‘Central’, ‘Neck’ and ‘Sleeve’ but no significant difference between groups was found for AOI ‘Shirt’ (while the AOI ‘Details’ was not present). Conversely, in training 2 the experimental group fixated significantly less than control in the AOI ‘Shirt’ and no difference was found in other AOIs beside ‘Sleeve’ where the control group looked longer. Further, results of the pre and post-test pointed out that the experimental group looked significantly longer in the AOI ‘Central’ than control. However, the other AOIs showed no significant difference between the groups.

Therefore, our hypothesis can be confirmed in the AOIs ‘Central’, ‘Neck’ and ‘Sleeve’ of training 1, ‘Shirt’ of training 2 and the AOI ‘Central’ of the post-test. However, these results indicate that the impact of annotations was limited to when present on screen in the training session and that, in their absence, the experimental group looked longer than control only in one AOIs at post-test. Also, if training 1 seems to have affected the apprentices’ gaze as expected that was not true for training 2 where, beside ‘Shirt’, no difference was found in the AOIs for the experimental group. This partially unexpected result might be traceable to the modality effect (Glenberg et al., 1989). This effect indicate that cueing performed with both visual and audio is more effective than cueing that uses only one of the two channels (see Mayer, 2002). In this case, cueing in training 1 was performed with annotations appearing one after another while matched by an audio description of each of the annotated area, while training 2 had no audio and annotations were presented all at the same time for the duration of the image on screen (40 s). Training 1 potentially induced the apprentices of the experimental group to focus only on the annotated areas present on screen in that moment and ignore the rest; while in training 2 they had more freedom to observe the different sections of the image as well as the control group who viewed a non-annotated image for the 40 s. Although these mixed results, annotations seem to still have an effect on the experimental group that observed significantly less non-annotated part of the shirt (AOI ‘Shirt’).

Are There Significant Differences in Terms of Gaze Coverage Between the Groups?

We expected the experimental group to have achieved at post-test a larger coverage of the image since annotations were placed in all the most relevant parts of the shirts to train the apprentices to observe certain areas but also to observe the garment in its totality. But a non-significant—merely statistically significant—difference was found between the two groups even if the experimental group covered a larger portion of the image at post-test compared to control group. This result indicates that annotations did not have the desired effect in attracting participants’ attention to specific areas but not across the whole image. However, this result is not completely out of the ordinary since fostering coverage of images can be a complex task and sometimes, after interventions, the increased coverage is limited (see Eder et al., 2020).

Are There Significant Differences in Terms of the Verbalisation Score Between the Groups?

In contrast to expectations, it was found that the experimental group did not mention more details than the control group (even if in the post-test, there was an increase in

their verbalisation score). Moreover, students were able to describe a high number of details even at pre-test, and this especially true for the control group; this might have caused the training to not be as impactful as expected. Further, the request to verbalise could also have influenced participants' gaze since concurrent think-aloud can affect the manner in which participants look at visual materials (e.g. Hertzum et al., 2009; Oh et al., 2013). Further, verbalisations were added to the study as a second measure for expertise and were introduced since the teachers indicated that students are required to describe images while they perform class activities.

Conclusions

In summary, our results indicate that annotations attract students' gaze in most of the AOIs in training 1 but in just one at post-test and training 2. Also, annotations did not affect the apprentice's gaze coverage or verbalisations as strongly as expected.

We should note some limitations that could have influenced the results. For instance, the task and stimuli were selected in agreement with the teachers to create an ecological task and to be profession-specific, as suggested by Jarodzka et al. (2017); however, the images might have been too simple to identify a huge difference between the two groups compared to images such as radiographies or technical drawings. Also, results from the training sessions indicate that the format of training 1 is more effective in attracting the gaze of the apprentices since visual annotations are presented one after another and paired with audio description of the garment, a strategy that has been proved useful in many studies (see Mayer, 2002). This could be exploited even further by positioning an high number of annotations on the images to recreate an 'ideal trajectory' for the participants or, alternatively, use model's eye movements to show participants the best way, according to the expert, to explore the images (see, Jarodzka et al., 2013). Further, concurrent think-aloud during pre and post-test and training 2 could have had contrasting impact with the annotations and, therefore, retrospective think-aloud should be considered as an option in future studies on cueing. Another potential limitation is the length of the two training sessions that amounted to only few minutes; a longer session or multiple sessions repeated during weeks or months could have led to better results; however, students were at the beginning of their second year and they would have started lessons on shirts very soon, thereby impacting the ecology of the study. It is possible that second-year students might be already familiar with the task (already trained in the first year although on different kind of garments, in particular skirts) to actually benefit from the training and that it might be more valuable for younger students.

Despite its limitations, this study participates in the discussion of numerous questions in vocational education. For instance, it constitutes one of the few examples exploring the initial training of fashion designers. It also contributes to the small but fundamental literature on visual expertise among fashion designers and strengthens its importance in the profession; being exposed to experts' way of looking at a garment by examining their annotations seems a promising technique to improve students' visual skills at work. Finally, the study is ecological and education-based

since it builds on a task that is already used in class to teach students how to observe garments with the use of annotations.

Findings could be valuable for teachers to create further learning scenarios (Cattaneo et al., 2015) and implement new tasks to convey a systematic and professional-specific way to observe garments in photos, patterns, and material forms using annotations. For instance, teachers could use annotations in a wide range of tasks in which they show students multiple photos of different parts of the same garment to increase their understanding of the whole item, or while teaching them how to read and create patterns or potentially apply them to video showing manufacturing procedures and tasks.

Acknowledgements We thank Prof. Jean-Luc Gurtner, Prof. Carlo Tomasetto for assistance with the analyses and Mr. Andrea Conforto, Dr. Stefano Tardini and Mr. Christian Milani for technical assistance.

Funding Open Access funding provided by EHB Eidgenössisches Hochschulinstitut für Berufsbildung, IFFP Institut fédéral des hautes études en formation professionnelle, IUFPF Istituto Universitario Federale per la Formazione Professionale. This work is supported by the Swiss State Secretariat of Education, Research and Innovation (SERI) [1315000851].

Open Access This article is licensed under a Creative Commons Attribution 4.0 International License, which permits use, sharing, adaptation, distribution and reproduction in any medium or format, as long as you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons licence, and indicate if changes were made. The images or other third party material in this article are included in the article's Creative Commons licence, unless indicated otherwise in a credit line to the material. If material is not included in the article's Creative Commons licence and your intended use is not permitted by statutory regulation or exceeds the permitted use, you will need to obtain permission directly from the copyright holder. To view a copy of this licence, visit <http://creativecommons.org/licenses/by/4.0/>.

References

- Baayen, R. H., Davidson, D. J., & Bates, D. M. (2008). Mixed-effects modeling with crossed random effects for subjects and items. *Journal of Memory and Language*, 59(4), 390–412. <https://doi.org/10.1016/j.jml.2007.12.005>
- Boucheix, J. M., & Lowe, R. K. (2010). An eye tracking comparison of external pointing annotations and internal continuous annotations in learning with complex animations. *Learning and Instruction*, 20(2), 123–135. <https://doi.org/10.1016/j.learninstruc.2009.02.015>
- Caruso, V., Cattaneo, A., & Gurtner, J. L. (2017). Creating technology-enhanced scenarios to promote observation skills of fashion-design students. *Form@ re-Open Journal per la formazione in rete*, 17(1), 4–17. <https://doi.org/10.13128/formare-20159>
- Caruso, V., Cattaneo, A., Gurtner, J. L., & Ainsworth, S. (2019). Professional vision in fashion design: Practices and views of teachers and learners. *Vocations and Learning*, 12(1), 47–65. <https://doi.org/10.1007/s12186-018-09216-7>
- Cassin, B., Rubin, M. L., & Solomon, S. (1984). *Dictionary of eye terminology*, 10. Triad Publishing Company.
- Cattaneo, A. Motta, E., & Gurtner, J. L. (2015). Evaluating a mobile and online system for apprentices' learning documentation in vocational education: Usability, effectiveness and satisfaction. *International Journal of Mobile and Blended Learning (IJMBL)*, 7(3), 40–58. <https://doi.org/10.4018/IJMBL.2015070103>
- Chamberlain, R., & Wagemans, J. (2015). Visual arts training is linked to flexible attention to local and global levels of visual stimuli. *Acta Psychologica*, 161, 185–197. <https://doi.org/10.1016/j.actpsy.2015.08.012>

- Coppi, A. E., Cattaneo, A., & Gurtner, J. L. (2019). Exploring visual languages across vocational professions. *International Journal for Research in Vocational Education and Training (IJRVET)*, 6(1), 68–96. <https://doi.org/10.13152/IJRVET.6.1.4>
- D'Innocenzo, G., Gonzalez, C. C., Williams, A. M., & Bishop, D. T. (2016). Looking to learn: The effects of visual guidance on observational learning of the golf swing. *PLoS ONE*, 11(5), e0155442. <https://doi.org/10.1371/journal.pone.0155442>
- De Koning, B. B., Tabbers, H. K., Rikers, R. M., & Paas, F. (2007). Attention cueing as a means to enhance learning from an animation. *Applied Cognitive Psychology: the Official Journal of the Society for Applied Research in Memory and Cognition*, 21(6), 731–746. <https://doi.org/10.1002/acp.1346>
- Dodd, B. J., & Antonenko, P. D. (2012). Use of signalling to integrate desktop virtual reality and online learning management systems. *Computers & Education*, 59(4), 1099–1108. <https://doi.org/10.1016/j.compedu.2012.05.016>
- Eder, T. F., Richter, J., Scheiter, K., Keutel, C., Castner, N., Kasneci, E., & Huetting, F. (2020). How to support dental students in reading radiographs: effects of a gaze-based compare-and-contrast intervention. *Advances in Health Sciences Education*, 1–23. <https://doi.org/10.1007/s10459-020-09975-w>
- State Secretariat for Education, Research and Innovation. (2018). *Ordinanza della SEFRI sulla formazione professionale di base Creatrice d'abbigliamento/Creatore d'abbigliamento con attestato federale di capacità (AFC) (Fashion Designer's Training Regulation)*. Retrieved from: <https://www.fedlex.admin.ch/eli/cc/2013/784/4>
- Glenberg, A. M., Mann, S., Altman, L., Forman, T., & Procise, S. (1989). Modality effects in the coding reproduction of rhythms. *Memory & Cognition*, 17(4), 373–383. <https://doi.org/10.3758/BF03202611>
- Goodwin, C. (1994). Professional vision. *American Anthropologist*, 96(3), 606–633. <https://doi.org/10.1525/aa.1994.96.3.02a00100>
- Grant, E. R., & Spivey, M. J. (2003). Eye movements and problem solving: Guiding attention guides thought. *Psychological Science*, 14(5), 462–466. <https://doi.org/10.1111/1467-9280.02454>
- Heiser, J., & Tversky, B. (2006). Arrows in comprehending and producing mechanical diagrams. *Cognitive Science*, 30(3), 581–592. https://doi.org/10.1207/s15516709cog0000_70
- Hertzum, M., Hansen, K. D., & Andersen, H. H. (2009). Scrutinising usability evaluation: Does thinking aloud affect behaviour and mental workload? *Behaviour & Information Technology*, 28(2), 165–181. <https://doi.org/10.1080/01449290701773842>
- Holmqvist, K., Nyström, M., Andersson, R., Dewhurst, R., Jarodzka, H., & Van de Weijer, J. (2011). *Eye tracking: A comprehensive guide to methods and measures*. Oxford University Press. <https://doi.org/10.1080/17470218.2015.1098709>
- Jaarsma, T., Jarodzka, H., Nap, M., van Merriënboer, J. J., & Boshuizen, H. P. (2014). Expertise under the microscope: Processing histopathological slides. *Medical Education*, 48(3), 292–300. <https://doi.org/10.1111/medu.12385>
- Jarodzka, H. M., & Boshuizen, H. P. A. (2017). Unboxing the black box of visual expertise in medicine. *Frontline Learning Research*, 5(3), 167–183. <https://doi.org/10.14786/flr.v5i3.332>
- Jarodzka, H., Holmqvist, K., & Gruber, H. (2017). Eye tracking in educational science: Theoretical frameworks and research agendas. *Journal of Eye Movement Research*, 10(1), 1–18. <https://doi.org/10.16910/jemr.10.1.3>
- Jarodzka, H., Scheiter, K., Gerjets, P., van Gog, T., & Dorr, M. (2010). How to convey perceptual skills by displaying experts' gaze data. In *Proceedings of the 31st Annual Conference of the Cognitive Science Society*, 2920–2925.
- 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. *Learning and Instruction*, 25, 62–70. <https://doi.org/10.1016/j.learninstruc.2012.11.004>
- Jasani, S. K., & Saks, N. S. (2013). Utilizing visual art to enhance the clinical observation skills of medical students. *Medical Teacher*, 35(7), e1327–e1331. <https://doi.org/10.3109/0142159X.2013.770131>
- Judd, C. M., Westfall, J., & Kenny, D. A. (2012). Treating stimuli as a random factor in social psychology: A new and comprehensive solution to a pervasive but largely ignored problem. *Journal of Personality and Social Psychology*, 103(1), 54. <https://doi.org/10.1037/a0028347>
- Kok, E. M., Abed, A., & Robben, S. G. (2017a). Does the use of a checklist help medical students in the detection of abnormalities on a chest radiograph? *Journal of Digital Imaging*, 30(6), 726–731. <https://doi.org/10.1007/s10278-017-9979-0>

- Kredel, R., Vater, C., Klostermann, A., & Hossner, E. J. (2017). Eye tracking technology and the dynamics of natural gaze behavior in sports: A systematic review of 40 years of research. *Frontiers in Psychology*, 8, 1845. <https://doi.org/10.3389/fpsyg.2017.01845>
- Krupinski, E. A., Chao, J., Hofmann-Wellenhof, R., Morrison, L., & Curiel-Lewandrowski, C. (2014). Understanding visual search patterns of dermatologists assessing pigmented skin lesions before and after online training. *Journal of Digital Imaging*, 27(6), 779–785. <https://doi.org/10.1007/s10278-014-9712-1>
- Krupinski, E.A., Tillack, A.A., Richter L., Henderson, J.T., Bhattacharyya, A.K., Scott, K.M., ... & Weinstein, R.S. (2006). Eye-movement study and human performance using telepathology virtual slides. Implications for medical education and differences with experience. *Human Pathology*, 37(12), 1543–1556. <https://doi.org/10.1016/j.humpath.2006.08.024>
- Mason, L., Pluchino, P., Tornatora, M. C., & Ariasi, N. (2013). An eye-tracking study of learning from science text with concrete and abstract illustrations. *The Journal of Experimental Education*, 81(3), 356–384. <https://doi.org/10.1080/00220973.2012.727885>
- Mayer, R. E. (2002). *Multimedia learning*. In *Psychology of learning and motivation* (Vol. 41, pp. 85–139). Academic Press. [https://doi.org/10.1016/S0079-7421\(02\)80005-6](https://doi.org/10.1016/S0079-7421(02)80005-6)
- Mitchell, J., Maratos, F. A., Giles, D., Taylor, N., Butterworth, A., & Sheffield, D. (2020). The visual search strategies underpinning effective observational analysis in the coaching of climbing movement. *Frontiers in Psychology*, 11, 1025. <https://doi.org/10.3389/fpsyg.2020.01025>
- Moreno, F. J., Reina, R., Luis, V., & Sabido, R. (2002). Visual search strategies in experienced and inexperienced gymnastic coaches. *Perceptual and Motor Skills*, 95(3), 901–902. <https://doi.org/10.2466/pms.2002.95.3.901>
- Moreno, F. J., Saavedra, J. M., Sabido, R., Luis, V., & Reina, R. (2006). Visual search strategies of experienced and non-experienced swimming coaches. *Perceptual and Motor Skills*, 103(3), 861–872. <https://doi.org/10.2466/pms.103.3.861-872>
- Naghshineh, S., Hafler, J. P., Miller, A. R., Blanco, M. A., Lipsitz, S. R., Dubroff, R. P., Khoshbin, S., & Katz, J. T. (2008). Formal art observation training improves medical students' visual diagnostic skills. *Journal of General Internal Medicine*, 23(7), 991–997. <https://doi.org/10.1007/s11606-008-0667-0>
- Naweed, A., & Balakrishnan, G. (2014). Understanding the visual skills and strategies of train drivers in the urban rail environment. *Work*, 47(3), 339–352. <https://doi.org/10.3233/WOR-131775>
- Oertel, C., Coppi, A., Olsen, J. K., Cattaneo, A., & Dillenbourg, P. (2019). On the use of gaze as a measure for performance in a visual exploration task. In *European Conference on Technology Enhanced Learning*, 386–395. Springer, Cham. https://doi.org/10.1007/978-3-030-29736-7_29
- Oh, K., Almarode, J. T., & Tai, R. H. (2013). An exploration of think-aloud protocols linked with eye-gaze tracking: Are they talking about what they are looking at. *Procedia-Social and Behavioral Sciences*, 93, 184–189. <https://doi.org/10.1016/j.sbspro.2013.09.175>
- O'Halloran, R. M., & Deale, C. S. (2006). Developing visual skills and powers of observation: A pilot study of photo interpretation. *Journal of Hospitality & Tourism Education*, 18(3), 31–43. <https://doi.org/10.1080/10963758.2006.10696862>
- Pang, E., Wong, M., Leung, C. H., & Coombes, J. (2019). Competencies for fresh graduates' success at work: Perspectives of employers. *Industry and Higher Education*, 33(1), 55–65. <https://doi.org/10.1177/0950422218792333>
- Ravesloot, C., van der Schaaf, M., Haaring, C., Kruitwagen, C., Beek, E., Ten Cate, O., & van Schaik, J. (2012). Construct validation of progress testing to measure knowledge and visual skills in radiology. *Medical Teacher*, 34(12), 1047–1055. <https://doi.org/10.3109/0142159X.2012.716177>
- Richter, J., Scheiter, K., & Eitel, A. (2016). Signalling text-picture relations in multimedia learning: A comprehensive meta-analysis. *Educational Research Review*, 17, 19–36. <https://doi.org/10.1037/edu0000220>
- Roach, V. A., Fraser, G. M., Kryklywy, J. H., Mitchell, D. G., & Wilson, T. D. (2019). Guiding low spatial ability individuals through visual cueing: The dual importance of where and when to look. *Anatomical Sciences Education*, 12(1), 32–42. <https://doi.org/10.1002/ase.1783>
- Scheiter, K., Eder, T., Richter, J., Hüttig, F., Keutel, C., (2019). Dental medical students' competencies for identifying anomalies in x-rays: When do they develop? (Conference paper). *18th Biennial Conference of the European Association for Research on Learning and Instruction (EARLI)*.
- Schneider, S., Beege, M., Nebel, S., & Rey, G. D. (2018). A meta-analysis of how signalling affects learning with media. *Educational Research Review*, 23, 1–24. <https://doi.org/10.1016/j.edurev.2017.11.001>

- Shapiro, J., Rucker, L., & Beck, J. (2006). Training the clinical eye and mind: Using the arts to develop medical students' observational and pattern recognition skills. *Medical Education*, 40(3), 263–268. <https://doi.org/10.1111/j.1365-2929.2006.02389.x>
- Sharma, K., Jermann, P., & Dillenbourg, P. (2015). Displaying Teacher's Gaze in a MOOC: Effects on Students' Video Navigation Patterns. *Design for Teaching and Learning in a Networked World*, 325–338. Springer. https://doi.org/10.1007/978-3-31924258-3_24
- Shin, D., & Park, S. (2019). 3D learning spaces and activities fostering users' learning, acceptance, and creativity. *Journal of Computing in Higher Education*, 31(1), 210–228. <https://doi.org/10.1007/s12528-019-09205-2>
- Srivastava, K. (2013). Emotional intelligence and organizational effectiveness. *Industrial Psychiatry Journal*, 22(2), 97. <https://doi.org/10.4103/0972-6748.132912>
- Stürmer, K., Könings, K. D., & Seidel, T. (2013). Declarative knowledge and professional vision in teacher education: Effect of courses in teaching and learning. *British Journal of Educational Psychology*, 83(3), 467–483. <https://doi.org/10.1111/j.2044-8279.2012.02075.x>
- Styhre, A., & Gluch, P. (2009). Creativity and its discontents: Professional ideology and creativity in architect work. *Creativity and Innovation Management*, 18(3), 224–233. <https://doi.org/10.1111/j.1467-8691.2009.00513.x>
- van der Gijp, A., Ravesloot, C. J., Jarodzka, H., Van der Schaaf, M. F., Van der Schaaf, I. C., van Schaik, J. P., & Ten Cate, T. J. (2017). How visual search relates to visual diagnostic performance: a narrative systematic review of eye-tracking research in radiology. *Advances in Health Sciences Education*, 22(3), 765–787. <https://doi.org/10.1007/s10459-016-9698-1>
- van Gog, T. (2014). The signaling (or cueing) principle in multimedia learning. In R. E. Mayer (Ed.), *The Cambridge handbook of multimedia learning*, *Cambridge Handbooks in Psychology*, 263–278. Cambridge University Press. <https://doi.org/10.1017/CBO9781139547369.014>
- Wolff, C. E., Jarodzka, H., van den Bogert, N., & Boshuizen, H. P. A. (2016). Teacher vision: expert and novice teachers' perception of problematic classroom management scenes. *Instructional Science*, 44(3), 243–265. <https://doi.org/10.1007/s11251-016-9367-z>
- Wood, G., Knapp, K. M., Rock, B., Cousens, C., Roobottom, C., & Wilson, M. R. (2013). Visual expertise in detecting and diagnosing skeletal fractures. *Skeletal Radiology*, 42(2), 165–172. <https://doi.org/10.1007/s00256-012-1503-5>

Publisher's Note Springer Nature remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.

Alessia Eletta Coppi is a PhD student at Swiss Federal Institute for Vocational Education and Training (SFIVET). Her main research interests focus on vocational education and training, new technologies and games, usability and cognitive enhancement.

Catharine Oertel is an Assistant Professor at TU Delft and is part of the Interactive Intelligence Group. Her main research fields are multi-modal modelling of conversational dynamics, conversational agents, multimodal interactions, social signal processing, human–robot interaction and social robotics.

Alberto Cattaneo is a Professor at the Swiss Federal Institute for Vocational Education and Training (SFIVET) where he leads the Research Field “Innovation in VET”. His main research fields concern the integration of ICT – and especially video-based technology – in vocational teaching and learning, reflective learning in VET, instructional design, multimedia learning, teachers' training.