



Castledown

 OPEN ACCESS

# Technology in Language Teaching & Learning

ISSN 2652-1687

<https://www.castledown.com/journals/tltl/>

*Technology in Language Teaching & Learning*, 7(3), 102566 (2025)

<https://doi.org/10.29140/tltl.v7n3.102566>

## Remaking Signs and Potentials for Learning in L2 Multimodal Composing



PATTARAMAS JANTASIN<sup>a</sup> 

MÜGE SATAR<sup>b</sup> 

<sup>a</sup>*Language Institute,  
Thammasat University, Thailand*  
[pattaramas.j@litu.tu.ac.th](mailto:pattaramas.j@litu.tu.ac.th)

<sup>b</sup>*Department of Applied Linguistics and  
Communication, Newcastle University, UK*  
[muge.satar@newcastle.ac.uk](mailto:muge.satar@newcastle.ac.uk)

### Abstract

Multimodal composing activities invite language learners to draw upon distinct potentials of diverse modes of representation to make and remake signs for communication. Despite increased interest in incorporating multimodal composing activities in language classrooms, we do not yet have an in-depth understanding of the process of remaking signs and its implications on the meaning being expressed and potentials for learning. To address this gap, we investigate multimodal composing activities of 25 undergraduate students in one EFL classroom in Thailand, who designed digital posters of an innovative product in small groups and delivered in-class presentations. Through social semiotic analysis of the posters and video recordings of the presentations by two groups, we examine specific changes in meaning representation (gains and losses) during multimodal composing. Our findings reveal that remaking signs across different modes shaped not only the meaning being expressed, but also what became available for students to learn and how. Our conclusion emphasizes the need for teachers to fully utilize modes available in all learning environments to best support students in achieving their communicative and learning objectives.

**Keywords:** multimodal composing, EFL learners, social semiotics, gains and losses, potentials for learning

---

**Copyright:** © 2025 Pattaramas Jantasins & Müge Satar. This is an open access article distributed under the terms of the Creative Commons Attribution Non-Commercial 4.0 International License, which permits unrestricted use, distribution, and reproduction in any medium, provided the original author and source are credited. **Data Availability Statement:** All relevant data are within this paper.

## Introduction

*Multimodal composing* refers to the act of creating *multimodal texts*— semiotic entities of any kind whether in “two, three, or four dimensions” (Kress, 2011, p. 207) which incorporate multiple modes, such as writing, speech, colour, images, and videos, to represent meaning (Belcher, 2017; Li & Storch, 2017). Multimodal composing involves making and remaking signs using different *modes*, which are “socially shaped and culturally given semiotic resource[s] for making meaning” (Kress, 2010, p. 79). Each mode has unique potentials and constraints and provides different routes to learning by shaping their attention to different aspects of the world (Kress & Bezemer, 2015); therefore, non-linguistic modes are not mere illustrations but play a central role in meaning representation alongside speech and writing (Jewitt, 2005). When modes are used together, they each contribute to the overall meaning (Jewitt et al., 2016; Kress, 2010). Thus, all modes need to be considered together for a full(er) understanding of the meaning being made.

In foreign language classrooms, *multimodal composing* is gaining popularity, with activities like creating PowerPoint slides, performing presentations, producing posters, or writing micro-blog posts. Yet, many teachers still prioritize written or spoken language, particularly in activities aimed at developing second language writing skills. While research in *multimodal composing* often focuses on how different modes can be used as means for language development, this development is frequently limited to two modes: writing and speech (e.g. Dzekoe, 2017; Kim & Belcher, 2020; Vandommele et al., 2017). To succeed in today’s highly digital communication environments, language learners need more than the mastery of speech and writing; they must also understand the power of various modes, the ways they convey meaning as an ensemble, and how to deploy different modes for communication in meaningful and impactful ways (Jewitt & Kress, 2003; van Leeuwen, 2017). Thus, a narrow focus on writing and speech overlooks the broader potential of multimodality to enrich communication and limits our understanding of learners’ meaningful engagement in contemporary communication, where these modes rarely appear alone but are typically accompanied by others (Valdés, 2017).

Therefore, it is crucial for language teachers to (re)consider current pedagogical practice if they aim to support students’ communication skills in all its forms. It is important to offer students a wider range of modes to generate greater meaning potentials, expand learning opportunities, and connect with out-of-school digital literacy that is now common in contemporary communication (Kress, 2015; Marissa & Hamid, 2022). A number of researchers have embraced the *multiliteracies* perspective, which highlights the development of learners’ ability to negotiate meaning across modes of representation and different contexts (Cope & Kalantzis, 2009), e.g., relational pedagogy (Kern, 2015), distributed collaboration (DePalma & Alexander, 2018), and a pedagogic metalanguage of transpositional grammar (Christoforou, 2025; Lim & Tan-Chia, 2022).

Despite growing interest in multimodal pedagogical design (Lim, Toh, & Nguyen, 2022), our understanding of how language learners remake signs is still limited. Multimodal composing involves the process of making and remaking signs across modes. While some studies have shown how multimodal composing helps learners to express their intended meanings through making signs with various modes and fosters new learning experiences (e.g., Canale, 2019; Chen, 2021; Yang, 2012), exploring the process of remaking signs is significant because when signs are remade, “there can never be (an assumption of) a perfect translation” (Bezemer & Kress, 2017, p. 526), inevitably having “consequences for learning, for knowing and shaping information and knowledge, for attending to and communicating about the world and our place in it” (Kress, 2010, p.5). Moreover, most research on L2 students’ sign-making focused on one type of multimodal text, i.e. digital storytelling (Li & Akoto, 2021), limiting insights into how sign remaking occurs across multimodal text types. Since learners encounter various multimodal text types in real-life communication, examining a wider range of texts

is essential to understand distinct potentials of modes that lead to more effective communication. To fill this gap, we address the following research questions:

- How are signs remade from a poster to an in-class presentation as part of a multimodal composing activity?
- What are the gains and losses in meaning representation and the potentials for learning about the product?

In this article, we analyse multimodal compositions (posters and in-class presentations) created by two groups of language learners. We examine how meaning is transformed as it moves between modes – specifically from the visual and textual elements of a poster to the spoken, gestural, and physical elements of a live presentation. We investigate the meaning(s) instantiated in the multimodal texts and how they shape potentials for learning when composing across modes.

### **Literature Review: Research on Multimodal Composing and Re-making of Signs in L2 Contexts**

The majority of recent studies has focused on the importance and benefits of bringing multimodal composing into language classrooms. For example, research found that multimodal composing can empower knowledge sharing and self-expression (Hughes & John, 2009; Shin & Cimasko, 2008; Shin et al., 2020), increase students' engagement in classroom learning (Hafner & Miller, 2011; Liang & Lim, 2021) and prepare students for digital communication (Belcher, 2017).

Many studies on L2 multimodal composing have adopted *a weak version of multimodality* (Grapin, 2019), privileging language and treating non-linguistic modes as merely a bridge for writing and speaking proficiency development. Dzekoe (2017), for example, investigated how multimodal composing activities helps ESL students notice linguistic and rhetorical features that need revisions in their written essays. Similarly, Vandommele et al., (2017) examined the effects of collaborative multimodal composing on the academic writing development of beginner Dutch learners. Kim and Blecher (2020) compared EFL learners' multimodal texts and their traditional essays for complexity and accuracy of language output. Within this perspective, learning opportunities are created by self-expression in various semi-otic systems, identification of what signs were not or cannot be expressed in the linguistic system, and re-articulating those signs through linguistic resources to improve writing. While reinforcing linguistic skills, this view risks undervaluing the communicative potential of other modes and limits opportunities for learners to develop broader multiliteracies and to engage in authentic, multimodal communication.

In contrast, fewer multimodal composing studies have adopted *a strong view of multimodality* (Grapin, 2019), which treats non-linguistic modes not merely as scaffolds or supports, but as legitimate semiotic tools with different potentials and limitations for making meaning (Jewitt, 2009; Kress, 2010). Unlike the weak view, the strong view positions all modes as equally valuable for constructing and communicating meaning. Although this view has not been widely embraced in language learning yet (Grapin, 2019; Yi et al., 2021), it has been strongly encouraged by many researchers in order to meet the evolving communication needs effectively (e.g. Blecher, 2017; Hafner, 2014; Li & Akoto, 2021; Smith, 2017; Yi et al., 2020). Within the strong multimodality view, learning opportunities arises through exploring new associations between signs, deducing new meanings, and learning to express them through all available semiotic systems by choosing apt resources to extend one's semiotic repertoire.

Much of this body of literature has examined the use of modes in students' multimodal texts, yet the findings consistently reveal a tendency to privilege linguistic modes over others. For example, Shin and Cimasko (2008) investigated how ESL university students used different modes in a

web-based project. They found that non-linguistic modes were mainly used to illustrate meaning already expressed in written language and sometimes to represent emotional dimensions that were difficult to express in words. Hafner (2014) examined EFL students' scientific documentaries. He demonstrated the ways in which learners' interest in designing documentaries were shaped by differences in the assumed audience, the sample documentaries viewed, and the preference for being distinct, resulting in the creative use of modes and resources to express meaning and identities. Chen (2021) investigated EFL students' perceptions and found that most favoured multimodal composing as it helped them to convey their ideas more clearly and tangibly. However, they still viewed language as "an indispensable mode" (Chen, 2021, p. 30). The findings of these studies suggest that the students' selection and integration of modes were primarily driven by social practices and traditional academic norms and discourses (Hafner, 2014; Shin & Cimasko, 2008). Moreover, while the students combined linguistic and non-linguistic modes, the persistent reliance on linguistic modes reflects the weak multimodal view that has long dominated language learning, despite increasing scholarly efforts to adopt the strong multimodal view. This underscores the urgent need to embrace a strong view of multimodality that enables learners to fully exploit the meaning-making potential of all semiotic modes.

One critical yet overlooked aspect within the strong multimodal view is the role of remaking signs within and across modes, i.e. *transformation and transduction* (Bezemer & Kress, 2016). These processes have received limited scholarly attentions (Kress, 2005) although they are common in everyday interaction (Jewitt et al., 2016) and multimodal composing (Belcher, 2017; Miller-Cochran, 2017). Nelson (2006), for instance, investigated transduction and transformation in the digital story creation processes of L2 writers. Drawing on multiple data sources, multimodal composing provided resources for powerful expression and opportunities to recognise new semiotic relationships. The results suggested that understanding modal affordances and the ability to recombine multimodal resources are critical for enacting powerful self-expression. Yang (2012) investigated how transduction and transformation helped learners achieve their intended meaning. For instance, intonation marks in a written script guided the learner to represent different emotions of the character in her voice narration. Lee et al. (2021) explored how EFL students created booklets presenting Taiwanese cultures. Their study revealed the students engaged in recursive reading-writing processes using multimodal resources, e.g. guidebooks, blogs, or movies. Through remaking signs, they deepened their understanding of cultural elements, audience needs, and semiotic choices in text design. Likewise, investigated transformative engagement amongst EFL trainee teachers while carrying out collaborative virtual exchange tasks, e.g. creating a Padlet board, and demonstrated how the participants achieve multimodal representation, inclusive interactional practices, and intercultural affinity through engaging with different meaning-making resources.

These studies highlight how the remaking of signs enable learners to express intended meanings and generate new learning experiences. However, none of them focused on how shifts across modes result in gains and losses in meaning, and more importantly, how these changes shape learning potentials. There is a need for research that not only maps the semiotic processes in multimodal composing but also critically examines their impact on learners' meaning making and learning potentials. Moreover, there has been a call for studies conducted "with differing composers, contexts, genres, and tools" (Smith et al., 2017, p. 274), especially with L2 students (Belcher, 2017) in the field of L2 writing (Li & Akoto, 2021) to gain insights into language learners' multimodal composing. We aim to fill this gap by examining transduction of signs between Thai L2 learners' digital posters and in-class presentations.

### Theoretical Framework

To examine how meaning is transformed across modes, and how multimodal texts instantiate meanings that shape learning opportunities, we draw on social semiotics (Bezemer & Kress, 2016) and multiliteracies (The New London Group, 2000).

## Social Semiotics: Signs, Remaking Signs, and Learning as Transformative Engagement

Social semiotics examines meaning within socio-cultural contexts, focusing on how people use and develop modes to “represent their understanding of the world” (Bezemer & Jewitt, 2009, p. 1). Central to this is the concept of a *sign*, in which meanings are represented and become visible (Selander, 2008). *A sign is motivated* and consists of the signified (meaning) and signifier (form); a sign maker always chooses modes and semiotic resources that are appropriate for their intended meaning (Kress, 2010). In multimodal composing, like the one presented in this paper, sign makers may use pre-existing resources (limited to what is available online or offline) but their choices are always motivated by the meanings they wish to convey.

When composing multimodally, learners construct signs to represent their ideas and choose modes based on their interests and aspects of the sign they perceive as criterial (Kress, 2010). For instance, when creating a digital story, they select visual representations for the characters, scenes they wish to portray, and dialogue that reflect their interpretation of the story. During these processes, some aspects of the story and characters are retained and *remade as new signs*, while others are deleted, resulting in unique representations and realities in each digital story.

Through multimodal composing, learners remake signs either within the same mode (*transformation*) or across different modes (*transduction*) (Bezemer & Kress, 2016). Examples of transformation include re-making signs in meeting minutes to a written report or the translation of a Thai novel into English; and examples of transduction may include drawing a picture to re-represent a poem, or creating a video based on a written story. Such processes create new signs remade that can be considered as *signs of learning* (ibid; Towndrow et al., 2013). In social semiotics, learners reshape signs, evoke new associations, and deduce signs (meanings) in ways they might not have known before. For example, Canale (2019) shows how a student re-articulated “I go to high school at 7:30 am” into a cartoon image of a girl running to school with a specific position of arms and feet, a bag in hand, and smoke by the feet. While writing only evokes *what* one does and *when*, image enables the learner to explore and express new associations between signs in relation to *how* (in this case: alone, on foot, fast). The combination of writing and image provide a fuller aspect of meaning, signifying ‘I rush to high school at 7.30 am’. Multimodal composing, thus, provides opportunities to express themselves more fully and “to learn in the process of making as [students] explore ways of expressing their ideas” (Lim & Tan-Chia, 2022, p.74).

However, when signs are remade, especially across modes, this will always bring about *gains and losses in meaning* that are represented (Bezemer & Kress, 2017), which not only affect the meaning being communicated, but also impact *potentials for learning* and knowledge construction (Kress, 2010, p. 5). *Potentials for learning* refer to possibilities offered by a text or environment that provide “the ground for learning and in that way may shape what learning is and how it may take place” (Bezemer & Kress, 2008, p.168). For instance, Bezemer and Kress (2017) illustrate that when imagery is used to represent a human heart, there is a significant gain in focus because the image provides a frozen representation of how one patient’s heart looks. However, there is a loss in movement (pulsation of a heart), which is crucial for surgical trainees’ learning. Thus, understanding how each mode differently shapes meaning and what to be learned is critical as it enables language teachers to strategically mobilize available modes across learning environments in ways that best support students’ learning objectives and communicative purposes.

## Multimodal Composing: Writing as Design

Within the multiliteracies perspective, multimodal composing is a dynamic process of meaning making, or ‘design’ (Kern, 2000; The New London Group, 1996). Writers draw on *available designs*,

i.e., all semiotic systems (e.g., gestures, writing, etc.) and schematic resources available in a community, and transform them through *designing*, i.e. “the process of shaping emergent meaning [that] involves re-presentation and recontextualization” (The New London Group, 2000, p. 22). The transformed available design (*the redesigned*) becomes a newly available design for future acts of meaning-making. In other words, writers draw on their knowledge, experiences or other resources that are available (*available designs*) and transform them to compose a multimodal text (*designing*): this results in the new resources available – *the redesigned*.

Accordingly, multimodal composing requires students to engage in transformative meaning making by choosing, designing, and re-designing modes and semiotic resources to best represent their intended meaning (Kress, 2010; The New London Group, 1996). Writing as design emphasises the ability to “understand the specific ways of configuring the world which different modes offer” (Hyland, 2009, p. 59) and how modes interact in combination. Such awareness enables students to compose and interpret multimodal texts effectively (Jewitt, 2006), and also to expand their semiotic repertoires in ways that speech and writing alone cannot support. Consistent with social semiotics, students’ subjective understandings of the world are transformed through designing (Cope & Kalantzis, 2009; Kress, 2003). Both view learners as “fully makers and remakers of signs and transformers of meaning” (Cope & Kalantzis, 2009, p. 175). Importantly, this challenges the weak multimodal view and foregrounds learners as active meaning-makers who can draw on diverse modes to achieve communicative goals. The framework advances the argument for moving beyond a narrow focus on speech and writing in multimodal composing toward multiple modes that better align with the dynamics of contemporary communication.

## Methods

This research adopts a qualitative case study design (Merriam & Tisdell, 2016) in a naturally-occurring setting (the participants’ regular classroom environment) to investigate two groups of Thai university EFL students engage in multimodal meaning-making. It focuses on *gains and losses in meanings* during transduction from a poster to an in-class presentation and *potentials for learning* about the product.

### Participants and Context

The participants were 25 first-year science and technology at a Thai university, enrolled in a foundation English course during the summer semester of 2019. These students, considered low-intermediate EFL learners, had IELTS scores below 4.5 (equivalent to B1 on the CEFR scale) and were placed in this course based primarily on their IELTS scores and the university’s English placement test scores. The course, graded on a pass/fail grading basis, aimed to improve English communication competence, focusing on reading, listening, speaking, and writing skills through 15 weekly 3-hour sessions. The learning activities were mainly designed based on an *active learning approach* with multimedia integration. For the final project, students created a digital poster advertising an innovative product using any software and presented it groups of 4-5 (see appendix). The posters were created outside of class in week 14, and the presentations took place in week 15, lasting 7–10 minutes. The project aimed to encourage students to use the target language structure (comparative and superlative adjectives) and creativity in communication through group work. The students received no specific training on multimodality or presentation skills.

After the presentations, the audience (other students) anonymously voted for the best presentation, which did not affect grades. Voting instructions were vague, with no detailed evaluation rubric. Success was determined by audience votes, following Kress’s (2010) idea that the audience defines communication success. Two groups were selected for in-depth analysis: Group Studio Box, which

received the most votes and effectively used multiple modes, and Group Anastasia, which received the fewest votes and relied mostly on speech. These two groups were different in terms of how they used the modes. The former one combined multiple modes coherently while the latter relied primarily on speech in their presentation.

The instructor also assessed the posters and presentations for 20% of the final grade based on language accuracy, cohesiveness, and comprehensibility. However, this rubric did not influence the audience voting and was not considered in case selection or data analysis.

### Data Collection and Analysis

The data were multimodal compositions created as part of the final project, which were 1) the students' digital posters, and 2) video recordings of their presentations (see Table 1). The data came from two cases: Group Studio Box (8 votes), and Group Anastasia (1 vote). The innovative product of the Studio Box group was a small light box with built-in LED lights that can be used for photographing small objects with better picture quality. The innovative product of Group Anastasia was a smart watch which can be used to text, make calls, and monitor physical activity. The presentation of Group Studio Box group lasted approximately 6 minutes, and that of Group Anastasia lasted 4 minutes.

**Table 1** *Overview of Data Collection*

| Group      | Multimodal texts                                                                    |                                                                                                                                             |
|------------|-------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------------------------------------------------------------|
|            | Poster                                                                              | Presentation                                                                                                                                |
| Studio Box |  | <p>Screenshot from the video-recorded presentation</p>  |
| Anastasia  |  | <p>Screenshot from the video-recorded presentation</p>  |

Ethical approval and informed consent were obtained before data collection. All student names were anonymised throughout.

The data were analysed using social semiotics (Bezemer & Kress, 2016), focusing on transduction processes. The first step of the analysis was to identify the signs (unit of analysis) that were transduced across the two multimodal texts: the poster and the in-class presentation. Firstly, posters and video recordings of the presentations were repeatedly viewed and transcribed using transcription symbols devised from Jefferson (1984) (see the Appendix) and with screenshots to provide a multimodal account of the data. The data were then coded for the signs made in the texts using NVivo 12 Pro. This process involved identification of modes used and meanings represented in each text. Signs that were identified included the product name, appearance, functions, promotions, and price. The coded signs were then compared across the two texts to identify the signs that were remade through transduction to explore gains and losses in meaning during the process of transduction. In this paper, two signs were purposively selected, describing: 1) product appearance, and 2) product use (see Table 2). These signs were selected because they were unique and stood out from the others (Bezemer & Jewitt, 2009, 2010): they were remade through transductions across a variety of modes, and also appeared to explain why the groups received different numbers of votes from the audience.

**Table 2** *Overview of the Transduced Signs*

| Focal signs           | Group      | Modes in the poster | Modes in the live presentation      |
|-----------------------|------------|---------------------|-------------------------------------|
| 1. Product appearance | Studio Box | Image               | 3D object                           |
|                       | Anastasia  | Image               | Speech                              |
| 2. Product use        | Studio Box | Image               | Speech, gestures, and 3D object     |
|                       | Anastasia  | Writing             | Speech, hand gestures, and eye gaze |

The last step was to carefully examine the selected instances of transduction. To identify gains and losses in meaning between the original signs made in the poster and the remade signs in the presentation, we attended to the meaning being made and identified semiotic changes in representation in terms of, for example, specificity, generality, and idealisation (Bezemer & Kress, 2008). As differences in meanings entail consequences for learning (Kress, 2010), we then explored how transduction shaped potentials for learning, i.e. what was to be learned and how the students learned it.

## Analysis

The analysis presents two instances of transduction from two cases – Group Studio Box and Group Anastasia. The first focal sign we explore in each case is related to product appearance and the second one concerns product use.

### Winning the Crowd: The Power of Semiotic Repertoires

#### *The Product Appearance: Transduction from Poster (Image) to Presentation (a 3d Object)*

Group Studio Box used an image in their poster to describe the appearance of their innovative product: a small light box with built-in LED lights for photography. In their in-class presentation, however, the product was represented using a 3D object. In the poster, the students used a cluster of three smaller

images to visually depict the product appearance (see Figure 1). The upper image showed how the product looked when in use. The other two images showed the product from the same angle in darkness and in light and with two background colours: black and white. In the presentation, the students used a prototype product, and this 3D object depicted the product appearance (see Figure 2).



**Figure 1** Image Showing Product Appearance in the Poster.



**Figure 2** 3D Object in the Presentation.

On the one hand, the transduction from the mode of image to 3D object resulted in a gain in *specificity* (Bezemer & Kress, 2008). The image represents the product from a single angle only, whereas the 3D object can be explored from all angles; it can also be moved and touched. Such affordances expand meaning potentials, allowing the learners to represent a fuller meaning of the product appearance. In the artefact, meanings related to the actual size, shape, dimensions, and material substance of the product are represented. None of these can be represented through the use of image. On the other hand, there was a loss in *generality and idealisation* (Bezemer & Kress, 2008). The use of image depicts an ideal picture of the product taken from a specific angle under specific light. It entails epistemological implications about how all the products should look like in general. However, the 3D artefact represents one specific product: the Studio Box that was brought into the classroom at that moment, which might not be in an ideal condition. Also, the transduction entailed a loss in *variation*. While the images in the poster show

variance in product appearance under two different lights (dark and light) and with two background colours (black and white), the 3D object only signifies its appearance under the specific light condition (unless the environmental conditions can be changed) and with a white background.

These gains and losses in meaning have implications for expressiveness and potentials for learning about the product. Firstly, in the image, the angle of the product fixes the meaning available for interpretation and what is to be learned about the product in terms of its appearance: size, shape, dimensions, material. However, the 3D object can be touched, moved and tilted in line with one's own interests. This provides opportunities to engage with the actual product, gain deeper understandings, and draw inferences about the product's size, dimensions, strength and durability. Moreover, the 3D object also provides opportunities to learn about other aspects of the product. For example, how is the box lit up? Is there a switch to turn on the light? Is it possible to adjust the light level? Therefore, it can expand meaning potentials and increase engagement in learning in a way that is different from the mode of image.

### *The Product Use: Transduction from Poster (Image) to in-Class Presentation (Speech, Gestures, and a 3D Object)*

Transduction was also observed related to how product use was represented in the poster (image) and the in-class presentation (a combination of speech, gestures and 3D object). In the poster in Figure 1, the image on top showed a shell placed inside the Studio Box being photographed using a mobile phone. In the presentation (Extract 1), one student first briefly explained how to take pictures with the product and then employed speech, gesture and the object to demonstrate the steps of use in detail. Re-articulation of the signs related to product use from poster (image) to presentation (speech, gestures, and 3D objects) entailed various gains and losses.

#### **Extract 1. A Student Demonstrates How to use the Product**

Line 01 **JJ:** The Studio box size 20 by 23 by 24 centimetres ((*pause*)) has right ((*pause*)) has a light LED to connect the USB which ((*pause*)) the connect USB to the power bank ((*grab a wire and connected it to the power bank* (Fig.3) and then the light switched on (Fig.4))

(Fig. 3)



(Fig. 4)



Line 02 **Audience:** Wow! ((*rising voice*))

Line 03 **JJ:** Studio box has background ((*pause*)) two background colours, white and black. ((*grab a watch from the chair* (Fig.5)) ((*pause*)) you ((*pause*)) when you want to take picture some mini object picture you just ((*pause*)) it is put in the Studio box ((*put the watch inside the Studio box* (Fig.6))).

(Fig. 5)



(Fig. 6)



Line 04 **JJ:** and start taking picture right away ((grab a mobile phone)) I will show how it works ((hand over the mic to her group member and take pictures of the watch (Fig. 7)))

Line 05 **JJ:** ((grab the mic back and show her mobile phone screen to the audience (Fig. 8))) Tada! ((rising voice)) ((smile)) I taking picture by camera smartphone. ((pause)) And I get picture that look good and save ((pause)) and save time. ((put the mobile phone back on the chair)) (Fig. 7)



First, when product use was described in the presentation there was, compared to the poster, a gain in *specificity*. In the upper image in the poster (see Figure 1), only a single part of the action of taking a picture using the product is depicted. Foregrounded in the image is a hand holding a smartphone with a shell picture on the screen. In the background, the Studio Box is visible with the same shell inside. The image does not offer any other details as to how to use the product for photography, such as whether (and if so, when) the product needs to be plugged in. In the presentation, a wider range of modes for making signs are available: speech, gestures, and the 3D object, allowing the students to demonstrate how to use the product in more detail. The size and apparatus (LED light, USB, power bank) (line 1), background colours and where to place the object (line 3), and the need to use a smartphone to take a good picture and to save time (line 5) are described in speech. How to connect the product to a power bank using a cable and switch on the light (see Figure 3 and Figure 4), placing a watch inside the Studio Box (see Figure 5 and Figure 6), and positioning oneself in front of the Studio Box, using a smartphone to take a picture, and having a satisfactory photograph (see Figure 7 and Figure 8) are demonstrated through embodied actions using gesture and object (along with head movement, body movement, and gaze shifts). In speech, we observe self-corrections and pauses which indicate difficulty in linguistic expression. Thus, it is likely that, with speech alone, product use may not have been conveyed as efficiently. The students' modal choices to depict product use also indicate that they found actions and the object to be an apt resource to make the signs related to product use.

Second, there was a loss in *generality and idealisation* in product use, similar to the description of product appearance. The image represented a general, idealised way to use the product without the presence of a specific actor/photographer, making the image become more general and less personal. Third, the transduction from image to the combination of modes in the presentation also entailed a loss in *permanence*. In the image, signs remain as long as the image exists and is clearly visible. Unlike the image, the meanings were expressed in a linear sequence of actions in the presentation. Each action, or meaning, was temporarily instantiated. Once the signs were made, no visible trace of them remained, unless the performance was recorded and made available to the audience.

In terms of expressiveness and potentials for learning about the product, in the presentation, a gain in specificity provided opportunities to better understand product use. The combination of speech, gestures, and 3D objects conveyed additional meanings that were different from those represented in the image, such as connection to a power source, the brightness of the light, the distance of the photographer from the product, as well as the photograph being taken at that specific time. Combining modes

clearly contributes to fuller aspects of the meaning being represented and expands the opportunities to learn about the product. Demonstration of product use through multiple modes successfully drew audience attention: Extract 1, Line 2 shows a verbal audience reaction (“Wow!”). This evidences audience satisfaction with the routes into learning available through multiple modes in the presentation.

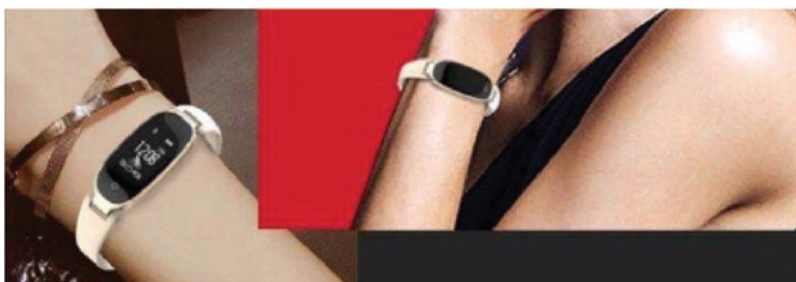
Yet, the presentation limited potentials for learning to some extent as the learning path was linear and fixed, with each action unfolding in a pre-determined sequence of action and information. In contrast, in the poster (see Figure 2), 3 images are displayed simultaneously, providing various entrances and exits for navigation and engagement. Therefore, the learning path is more flexible. The students, for instance, can design their own reading and learning route according to their interest. When the learning path is less rigid, learners have more freedom to establish new connections for themselves.

Finally, as the image has permanence, it also offers potential opportunities to engage with and make sense of the image multiple times. The presentation, however, needs to be recorded to enable multiple opportunities for learning at various times.

### Disjointed Multimodal Communication and Its Impact on Meaning-Making

#### *Describing Product Appearance: Transduction from Poster (Image) to Presentation (Speech)*

The innovative product of Group Anastasia was a smart watch which can be used to text, make calls, and monitor physical activity. In the poster, product appearance was represented in the mode of image (see Figure 9) showing the watch in two conditions: activated screen and deactivated screen from different angles. The front face of the digital watch with oval-shaped black screen and a beige thin strap is visible. In the image on the left, the activated watch screen shows time and weather conditions in white. Other functionalities cannot be identified due to image quality. In the presentation (see Extract 2), product appearance was depicted in the mode of speech describing it as a beautiful, compact, and affordable watch with a push button to make calls and text.



**Figure 9** Image Showing The Product Appearance in the Poster.

#### **Extract 2. The Students Introduce the Product in the Presentation**

**MK:** ((standing and talking (Fig. 10))) Hello everyone. Today we will talk about Anastasia. This product has a beautiful. ((pause)) It is com. It is compact and (nastly) and the price are affordable. This product has can take the (push) and talk on the phone and chat.



(Fig. 10)

The presentation entailed a gain in *attitude* and a loss in *specificity*. In Extract 2, depiction of the product as beautiful, compact, and affordable reflects the students' attitude towards the product's design, but it does not specify any aspects of the product appearance, such as shape, size, and colour, signs which are lost when transduced from image in the poster to speech in the presentation. Speech does not even specify the product as a watch, and it is also not clear which aspects of the product appearance make it beautiful. While the digital poster is displayed in the background during the presentation, image size and quality were poor (see Figure 10). The students stood still, without using gesture, and did not refer to the images on the poster which showed product appearance. Therefore, the vague images of a watch present in the background on the digital poster (which were not incorporated into the presentation) were the only signifiers of what product Anastasia was or what it looked like. If the presentation is treated separately, all details of product appearance are lost.

In terms of potentials for learning, the presentation, without reference to the poster, lacked sufficient information for new associations to be made about the product's specific appearance. Each learner could potentially imagine a very different design from the original product. The signs of attitude (beautiful, compact, and affordable) are only expressed in speech, which is the single inroad (mode) into learning (new associations).

### ***Describing Product use: Transduction from Poster (Writing) to Presentation (Speech, Hand Gestures, and Eye Gaze)***

We now explore transduction of signs related to product use, i.e. 'how to turn on the watch screen' from the poster (in the mode of writing) to the presentation (in the modes of speech, hand gestures, and eye gaze). In the poster, writing describes how to turn on the product screen (see Figure 11 left, and Table 1). Writing centred and placed on top of the poster stating "Wake up the screen with a lift of arm. Simple and cool" is a signifier that the screen can be turned on by lifting an arm and reflects *attitude* towards this functionality: it is "simple and cool". It occupies about top one third of the total space. The top position of writing, capital letters, font size and colours indicate the designers' interest in making the sign salient. The rest of the poster (bottom two-thirds) shows a woman wearing the watch on the right and writing, followed by a picture of the watch on the left (see Table 1). The writing on black background in white font reads "ANASTASIA" in italicised capitals and a large font, followed by "Turn your arm and the screen will be awakeny" in smaller font and sentence case (see Figure 11, right). Despite a grammatical mistake, this text conveys the concept of how to turn on the screen, but it no longer reflects *attitude* (sign-makers' opinion). The use of the conjunction 'and' and the (mistakenly written) future simple 'will be' indexes a sequence of events, indicating one event (turn your arm) happens and then another (the screen will wake up) happens.

In the presentation (see Extract 3), speech was the primary mode used. The presenter's word choice, word order, and position of pauses (often placed before a verb) framed her speech into many chunks that were different from the expected SVO order in English sentences. These influence the meaning being made, making it difficult to comprehend. For example, in Extract 3, line 02 and line 03, "and

**WAKE UP THE SCREEN WITH A LIFT  
OF ARM. SIMPLE AND COOL.**

Turn your arm and  
the screen will be awakeny

**Figure 11** *Writing in the Poster Signified how to turn on the Screen.*

gently (pause) make a turn the screen (pause) will be a (pause) will be awaken” appears to signify a cause-effect relation of two actions: gently turn the watch and then its screen will wake up. Yet, the writing on the poster suggests a slightly different sign: turn your arm and the screen will wake up.

### Extract 3. A Student Explains how to turn on the Watch Screen

Line 01 **SR:** How to use [wear] on the [wrist] ((*pause*)) [rise your arm.]

[(*looking up*)] (Fig. 12)

[(*turning to and looking at the poster*)] (Fig. 14)

[(*looking up*)] (Fig. 13)

(Fig. 12)



(Fig. 13)



(Fig. 14)



Line 02 **SR:** And [gently] ((*pause*)) make a turn the [screen] ((*pause*))

[(*raised her arm up with clenched fist*)] (Fig. 15)

[(*looking up*)] (Fig. 15)

[(*turning to and looking at the poster*)] (Fig. 16)

(Fig. 15)



(Fig. 16)



Line 03: **SR:** will be [a ((*pause*))] [will be awaken]]. No need to use another arm to check dates.

[(*looking up*)] (Fig. 17)

[(*turning to and looking at her group member*)] (Fig. 18)

[(*raised her arm up with clenched fist*)] (Fig. 17)

(Fig. 17)



(Fig. 18)



While primarily relying on speech to describe product use in the presentation, other modes, such as eye gaze and hand gestures, were also used. Combining various modes tends to expand meaning potentials and can lead to a gain in *specificity*, as in the analysis presented in Section 4.1.2 from Group Studio Box. However, this was not the case here. Speech, eye gaze, and hand gestures were not combined coherently and did not contribute to the meaning being made about product use. Instead, the way these modes were used and appeared together likely caused confusion in meaning and led to a loss of *expressiveness*. During the presentation describing the product use, the presenter averted her gaze away from the audience and upwards, and then turned her body to look at the poster and another group member (see Figures 12–18). She also raised her arm up with a clenched fist twice during the pauses while gazing upwards. Gaze shifts and gestures were not employed to refer to relevant writing on the

poster and the gestures did not index turning the arm to wake up the screen (with a watch on the arm and through wrist movements). Instead, several pauses in speech, eye gaze direction, and gestures together seemed to reflect attempts at recalling, potentially, a script prepared beforehand for the presentation. Thus, although the affordances of these modes can enable further meanings, the presenter did not utilise them to either complement the meaning being made or to provide additional detail as to how to turn on the screen.

Regarding potentials for learning, transduction of the signs related to product use from the poster to the presentation did involve employment of multiple modes, which would be expected to invoke multiple inroads into learning. Yet, since speech, eye gaze, and gesture shifts did not corroborate or complement the sign being made ('turn your arm to activate the screen'), they seemed to interfere with the meaning being expressed, making it more difficult to represent how to turn on the screen in the presentation compared to writing in the poster. When multiple modes were not orchestrated and were instead used randomly and unintentionally, the combination of these modes drew the attention away from the meaning being delivered and thus limited opportunities for learning and making new associations.

## Discussion

While few studies adopting a strong view of multimodality seek to understand how modes and semi-otic resources are used and developed for meaning making in students' multimodal texts (e.g. Hafner, 2014; Shin & Cimasko, 2008), studies on the remaking of signs are scarce (Kress, 2005), and further research in multimodal composing in different contexts is needed (Blecher, 2017; Smith et al., 2017). This study, therefore, examined the multimodal texts of Thai EFL students to gain insights into the remaking of signs through transduction, its implications for the meanings represented, and, in turn, the potentials for learning. Within a social semiotic perspective, every mode has distinct meaning potentials and "offers its own distinct route to learning" (Bezemer & Kress, 2016, p.61). Gaining insights into transduction is crucial as all transductions from one mode to the other inevitably entails gains and losses in meaning, which results in differences in expressiveness and potential opportunities for learning.

Likewise, the analyses in this paper revealed that when meanings in images were remade in speech, gestures, and a 3D object, various gains and losses were observed, which shaped what was to be learned. In the three instances of transduction (see Sections 4.1.1, 4.1.2, and 4.2.2), there was a gain in *specificity*. The 3D object used by Group Studio Box afforded touch and movement, while image did not. The 3D object, therefore, enabled potential additional connections such as the texture, or material of the product, its strength, and durability.

Moreover, the analyses showed that when a wider range of modes were employed, potentials for meaning making and for learning expanded. In Section 4.1.2, Extract 1, while the student exhibited signs of difficulty in her linguistic expression, she was still able to orchestrate her speech with gestures and a 3D object to communicate her message successfully to the audience. The student showed her ability to creatively bring together available modes and resources to make her intended meanings visible. As Bezemer and Kress (2016) argue, limiting students to rely on only one mode to represent meaning or knowledge would restrict resources to evidence what they know or have learned. Similarly, Shin and Cimasko (2008) and Hughes and John (2009) found that the use of multiple modes can empower self-expression and knowledge sharing. In a similar vein, Jewitt (2009), Bezemer and Kress (2016), Grapin, (2019), and Canale (2019) strongly argue for the expansion of a range of modes in learning environments to increase meaning potentials and widen opportunities for learning.

However, the ability to effectively utilise multiple modes to successfully convey meanings does not appear to exist naturally. The analysis of the last instance of transduction (see Section 4.2.2) shows a loss in expression due to an ineffective combination of multiple modes. This reminds us that even though students nowadays are exposed to multimodal texts and multimodal meaning making in their everyday lives, we cannot assume that they know how to best use available modes to make meanings in multimodal texts. This is in line with van Leeuwen (2017), Hyland (2009), and Jewitt and Kress (2003), who all argue that students need to understand the power of modes, the different affordances they bring, and how they work together to be able to employ them in meaningful and impactful ways. Kern (2015) recommends a relational pedagogy in which language educators are strongly advised to take this issue into account if they want to help students improve communication skills in its full form in contemporary communication practices.

Overall, the analyses illustrated how employment of different modes affect meaning representation, which entailed significant implications for what does and does not become available for students to learn, as well as in which order learning potentials become available. Our findings provided further evidence that each mode has unique possibilities for meaning making and draws attention to different aspects and understanding of the world (Bezemer & Kress 2016; Kress, 2010). The findings align with Nelson (2006) and Bezemer & Kress (2017), emphasizing the impact of modal choices on learning. In the current demands of multimodal and digital communication, the findings show the need to implement multimodal composing in L2 learning and teaching not only to strengthen the learners' speech and/or writing (weak view of multimodality) but to provide them opportunities to expand their semiotic repertoires for effective meaning making (strong view of multimodality). This study contributes to the existing literature by demonstrating how multimodal composing studies can focus on multiliteracies development. This study thus highlights the importance of educators expanding their focus beyond language to the interplay of modes, understanding their differences, and leveraging their combinations to enhance communication and shape learning outcomes.

## Conclusion

This study investigated how the signs made in posters were remade in different modes in in-class presentations of EFL students at a university in Thailand to explore what meanings were gained and lost through the process of transduction, and the implications of this for expressiveness and potentials for learning. Since we live in a world where multiple modes are commonly used together for making signs, understanding how different modes can be used, and how they draw attention to and shape meaning and learning in different ways are more important than ever before. With such understandings, teachers can design environments comprising modes conducive to learning objectives and encourage students' development of multimodal composing practices.

This article was a response to Smith et al.'s (2017) call for more multimodal composing studies in different contexts and Kress's (2005) call for studies focusing on the remaking of signs and potentials for learning. This study was conducted with a small data set and thus the generalisability of the findings may be limited. Data were group products, where selected modes and resources to be used in the outputs could have been influenced by each individual's interest, common practices in students' out-of-school life, and assignment genres. These topics were not examined in this study, but they could be addressed in future research.

This study makes significant contributions to the field of TESOL, second language writing, and social semiotics by providing insights into the agency of students as active sign-makers and their transformative meaning making processes when producing two different kinds of multimodal texts in an EFL context (more specifically, in the remaking of signs across modes as an inroad into understanding

learning). Investigating how modal choices shape meaning representation and potentials for learning is important. Not only can it help students expand their meaning-making potentials and better express themselves through a variety of modes, but also help teachers and educators design learning materials and environments that can shape learners' engagement and bring them closer to meeting set teaching objectives.

### Acknowledgements

For the purpose of open access, the author has applied a Creative Commons Attribution (CC BY) licence to any Author Accepted Manuscript version arising from this submission.

### References

- Belcher, D. (2017). On becoming facilitators of multimodal composing and digital design. *Journal of Second Language Writing*, 38, 80–85. <https://doi.org/10.1016/j.jslw.2017.10.004>
- Bezemer, J., & Jewitt, C. (2009). Social Semiotics. In *Handbook of Pragmatics* (pp. 1–13). John Benjamins Publishing Company. <https://doi.org/10.1075/hop.13.soc5>
- Bezemer, J., & Jewitt, C. (2010). Multimodal Analysis: Key issues. In L. Litosseliti (Ed.), *Research Methods in Linguistics* (pp. 180–197). Continuum International Publishing Group. <https://doi.org/10.1075/hop.13.soc5>
- Bezemer, J., & Kress, G. (2008). Writing in multimodal texts: A social semiotic account of designs for learning. *Written Communication*, 25(2), 166–195. <https://doi.org/10.1177/0741088307313177>
- Bezemer, J., & Kress, G. (2016). *Multimodality, learning and communication: A social semiotic frame*. Routledge.
- Bezemer, J., & Kress, G. (2017). Continuity and change: Semiotic relations across multimodal texts in surgical education. *Text & Talk*, 37(4), 509. <https://doi.org/10.1515/text-2017-0014>
- Canale, G. (2019). Toward a multimodal socio-semiotic account of learning. In *Technology, Multimodality and Learning: Analyzing meaning across scales* (pp. 41–82). Palgrave Macmillan. [https://doi.org/10.1007/978-3-030-21795-2\\_3](https://doi.org/10.1007/978-3-030-21795-2_3)
- Chen, C. W.-Y. (2021). Composing print essays versus composing across modes: Students' multimodal choices and overall preferences. *Literacy*, 55(1), 25–38. <https://doi.org/10.1111/lit.12234>
- Christoforou, M. (2025). Gen AI-assisted multimodal meaning Design: exercising a pedagogic meta-language of transposition. *Pedagogies: An International Journal*, 1–25. <https://doi.org/10.1080/01554480X.2025.2522884>
- Cope, B., & Kalantzis, M. (2009). “Multiliteracies”: New literacies, new learning. *Pedagogies: An International Journal*, 4(3), 164–195. <https://doi.org/10.1080/15544800903076044>
- DePalma, M.-J., & Alexander, K. P. (2018). Harnessing writers' potential through distributed collaboration: A pedagogical approach for supporting student learning in multimodal composition. *System*, 77, 39–49. <https://doi.org/10.1016/j.system.2018.01.007>
- Dzekoe, R. (2017). Computer-based multimodal composing activities, self-revision, and L2 acquisition through writing. *Language Learning & Technology*, 21(2), 73–95. <https://doi.org/10.125/44612>
- Grapin, S. (2019). Multimodality in the new content standards era: Implications for English learners. *TESOL Quarterly*, 53(1), 30–55. <https://doi.org/10.1002/tesq.443>
- Hafner, C. A. (2014). Embedding Digital Literacies in English Language Teaching: Students' Digital Video Projects as Multimodal Ensembles. *TESOL Quarterly*, 48(4), 655–685. <https://doi.org/10.1002/tesq.138>
- Hyland, K. (2009). *Academic discourse: English in a global context* (1 ed.). Bloomsbury Academic. <https://doi.org/10.5040/9781474211673>

- Jefferson, G. (1984). Notes on a systematic deployment of the acknowledgement tokens “Yeah”; and “Mm Hm”. *Paper in Linguistics*, 17(2), 197–216. <https://doi.org/10.1080/08351818409389201>
- Jewitt, C. (2005). Multimodality, “Reading”, and “Writing” for the 21st Century. *Discourse: Studies in the Cultural Politics of Education*, 26(3), 315–331. <https://doi.org/10.1080/01596300500200011>
- Jewitt, C. (2006). *Technology, literacy and learning: A multimodal approach*. Routledge.
- Jewitt, C. (2009). *Routledge handbook of multimodal analysis*. Routledge.
- Jewitt, C., Bezemer, J., & O’Halloran, K. (2016). *Introducing multimodality*. Routledge.
- Jewitt, C., & Kress, G. (2003). *Multimodal literacy*. Peter Lang.
- Kern, R. (2000). *Literacy and language teaching*. Oxford University Press.
- Kern, R. (2015). *Language, literacy, and technology*. Cambridge University Press. <https://doi.org/10.1017/CBO9781139567701>
- Kim, Y., & Belcher, D. (2020). Multimodal composing and traditional essays: Linguistic performance and learner perceptions. *RELC Journal*, 51(1), 86–100. <https://doi.org/10.1177/0033688220906943>
- Kress, G. (2005). Gains and losses: New forms of texts, knowledge, and learning. *Computers and Composition*, 22, 5–22. <https://doi.org/10.1016/j.compcom.2004.12.004>
- Kress, G. (2010). *Multimodality: A social semiotic approach to contemporary communication*. Routledge.
- Kress, G. (2011). Discourse analysis and education: A multimodal social semiotic approach. In R. Rogers (Ed.), *An introduction to critical discourse analysis in education* (pp. 205–226). Routledge.
- Kress, G. (2015). Semiotic work: Applied Linguistics and a social semiotic account of Multimodality. *AILA Review*, 28(1), 49–71. <https://doi.org/10.1075/aila.28.03kre>
- Kress, G., & Bezemer, J. (2015). A social semiotic multimodal approach to learning. In D. Scott & E. Hargreaves (Eds.), *The SAGE handbook of learning*. SAGE Publications Ltd. <https://doi.org/10.4135/9781473915213>
- Lee, S.-Y., Gloria, L. Y.-H., & Chin, T.-C. (2021). Practicing multiliteracies to enhance EFL learners’ meaning making process and language development: A multimodal Problem-based approach. *Computer Assisted Language Learning*, 34(1–2), 66–91. <https://doi.org/10.1080/09588221.2019.1614959>
- Li, M., & Akoto, M. (2021). Review of recent research on L2 digital multimodal composing. *International Journal of Computer-Assisted Language Learning and Teaching (IJCALLT)*, 11(3), 1–16. <https://doi.org/10.4018/IJCALLT.2021070101>
- Li, M., & Storch, N. (2017). Second language writing in the age of CMC: Affordances, multimodality, and collaboration. *Journal of Second Language Writing*, 36, 1–5. <https://doi.org/10.1016/j.jslw.2017.05.012>
- Lim, F. V., & Tan-Chia, L. (2022). *Designing learning for multimodal literacy: Teaching viewing and representing*. Routledge. <https://doi.org/10.4324/9781003258513>
- Marissa, N. D., & Hamid, S. (2022). Bridging out-of-school digital literacy through multimodal composition for EFL students with developing proficiency. *The JALT CALL Journal*, 18(2).
- Merriam, B. S., & Tisdell, J. E. (2016). *Qualitative research: A guide to design and implementation* (4th ed.). John Wiley & Sons.
- Miller-Cochran, S. (2017). Understanding multimodal composing in an L2 writing context. *Journal of Second Language Writing*, 38, 88–89. <https://doi.org/10.1016/j.jslw.2017.10.009>
- Nelson, M. E. (2006). Mode, meaning, and synaesthesia in multimedia L2 writing. *Language Learning & Technology*, 10(2), 56–76.
- Satar, M., Hauck, M., & Bilki, Z. (2023). Multimodal representation in virtual exchange: A social semiotic approach to critical digital literacy. *Language Learning & Technology*, 27(2).
- Selander, S. (2008). Designs for learning - A theoretical perspective. *Designs for Learning*, 1, 4. <https://doi.org/10.16993/dfl.5>

- Shin, D.-s., & Cimasko, T. (2008). Multimodal composition in a college ESL class: New tools, traditional norms. *Computers and Composition*, 25(4), 376–395. <https://doi.org/10.1016/j.compcom.2008.07.001>
- Smith, B., Pacheco, M., & Rossato de Almeida, C. (2017). Multimodal codemeshing: Bilingual adolescents' processes composing across modes and languages. *Journal of Second Language Writing*, 36, 6–22. <https://doi.org/10.1016/j.jslw.2017.04.001>
- The New London Group. (1996). A Pedagogy of multiliteracies: Designing social futures. *Harvard Educational Review*, 66(1), 60–93. <https://doi.org/10.17763/haer.66.1.17370n67v22j160u>
- The New London Group. (2000). A pedagogy of multiliteracies designing social futures. In B. Cope & M. Kalantzis (Eds.), *Multiliteracies: Literacy learning and the design of social futures* (pp. 9–37). Routledge.
- Valdés, G. (2017). Entry Visa Denied: The construction of symbolic language borders in educational settings. In G. Ofelia, F. Nelson, & S. Massimiliano (Eds.), *The Oxford Handbook of Language and Society*. <https://doi.org/10.1093/oxfordhb/9780190212896.013.24>
- van Leeuwen, T. (2017). Multimodal literacy. *Viden om Literacy*, 21, 4–11. [https://www.videnom-laesning.dk/media/2127/21\\_theo-van-leeuwen.pdf](https://www.videnom-laesning.dk/media/2127/21_theo-van-leeuwen.pdf)
- Vandommele, G., Van den Branden, K., Van Gorp, K., & De Maeyer, S. (2017). In-school and out-of-school multimodal writing as an L2 writing resource for beginner learners of Dutch. *Journal of Second Language Writing*, 36, 23–36. <https://doi.org/10.1016/j.jslw.2017.05.010>
- Yi, Y., Shin, D.-S., Cimasko, T., & Chen, K. (2021). Methodological approaches to examining multimodal composing. In D.-S. Shin, T. Cimasko, & Y. Yi (Eds.), *Multimodal Composing in K-16 ESL and EFL education: Multilingual perspectives* (pp. 17–31). Springer Singapore. [https://doi.org/10.1007/978-981-16-0530-7\\_2](https://doi.org/10.1007/978-981-16-0530-7_2)

## Appendix

### An ‘innovative product’ task taken from TU050 coursebook

#### Coursebook content

#### 6.6 Writing: Our product is the best!

**Scenario:** You are in a marketing team and you are told to prepare an ad for promoting your newest company product in the upcoming World Tech Expo—the world’s leading event which will introduce and explore the latest innovations.



**Instructions:** In groups, choose a product, perhaps something that you have bought recently or imagine a new one. Brainstorm possible adjectives for promoting your product. Name your product and create an advertisement on the given paper. Then be ready to present to the class and vote for the best ad!

Use comparative and superlative adjectives to compare your product with others and show that your product is better or has improved.



#### Handout created by the lecturer to support coursebook instructions

#### Final Project: Our product is the best!

This assignment is a group activity (4–5 people). Each group needs to create an advertisement to promote your newest company product and present it to customers at the World Tech Expo—the world’s leading event for introducing and exploring the latest innovations.

#### Task:

1. Each group chooses one product, perhaps something that you have bought recently or imagine a new one.
2. Name the product and brainstorm how to make your product better and more attractive. The students are encouraged to research current trends and be creative.
3. Create an A4 advertisement (paper or electronic file) and **use comparative and superlative adjectives** to compare your product to that of a competitor and to show that your product is better or has improved.
4. Prepare to perform a 10-minute presentation in class. The presentation content should include: 1) background and rationale; 2) description of the products (explain why your product is the best) (you can compare your product to the others in the market.); and 3) target customer, price, and promotion (optional).

**NOTE:** The lecturer might consider holding a competition and vote for the best advertisement. The winner might be awarded a prize.

**Transcription symbols devised from Jefferson (1984)**

| Symbol | Definition and Use              |
|--------|---------------------------------|
| [ ]    | Overlapping talk or action      |
| (word) | Uncertain word                  |
| (( ))  | Analyst comment or descriptions |
| !      | Rise in intonation              |