

Facilitating Interaction-Oriented AI Literacy in Translator Training: A Process-Oriented Approach

Erik Angelone

Institute of Translation and Multilingual Communication
TH Köln – University of Applied Sciences Cologne, Germany
erik.angelone@th-koeln.de

Abstract

This paper proposes a process-oriented approach using screen recording and think-aloud output for facilitating the AI literacy of translation students in the domain of interaction with generative AI. It outlines potential indicators of intelligence augmentation and impairment, as documented in process protocols of student performance, which students can use as a point of departure for self-reflecting on how generative AI utilization impacts their performance in terms of quality and efficiency.

1 Introduction

As generative AI continues to redefine translation workflows, the translator training community is responding with a range of learning activities to familiarize students with this assistive technology (see Penet et al., 2026; Pym and Hao, 2025). Translation-specific AI literacy (Krüger, 2005) is now being prioritized, informing competence frameworks and updates in translator training curricula. Thanks to recent surveys, we have a better understanding of *why* translators use generative AI. However, there is relatively little empirical research on *how* the utilization of generative AI impacts student translation performance in both positive and negative ways, from the perspective of quality and efficiency. This paper proposes a process-oriented approach that uses screen recording and think-aloud output to enable student self-reflection on their own performance along these lines when using generative AI. This approach is intended to foster

AI literacy in the domain of interaction, where metacognitive awareness plays a pivotal role.

2 Conceptual framework

Within the interaction domain of his translation-specific AI literacy framework, Krüger (2025, 18) foregrounds a cognitive level, with corresponding effects traced to interaction with generative AI. Conducive cognitive effects, such as efficiency gains and beneficial cognitive offloading, are regarded as forms of intelligence augmentation. Detrimental cognitive effects, such as AI priming and various forms of deskilling, are regarded as forms of intelligence impairment.

An inherent danger of deskilling lies in the fact that is often perceptual in nature (Appiah, 2025), and might go unnoticed. The same could hold true for perception of intelligence augmentation. In the context of process-oriented translator training, a combination of screen recording and think-aloud output enables the creation of translation process protocols documenting problems and corresponding problem-solving. These protocols can then serve as pedagogical framework for retrospective self-analysis, with students and instructors reflecting on salient indicators of performance. This approach can readily extend to explorations of intelligence augmentation and impairment in conjunction with generative AI interaction.

3 Pedagogical design and methodology

Initial instructor input on how to create a screen recording and the task of thinking aloud would scaffold the approach, which is otherwise largely rooted in self-regulated, discovery-based learning. If training is deliberately deductive, central ideas

pertaining to intelligence augmentation and impairment could be defined and contextualized upfront so that students know what to look for when reflecting on their performance. This starts with working definitions of each overarching concept and its respective forms. The instructor may also consider introducing concrete manifestations of augmentation and impairment found in each modality. Pedagogically speaking, it would be advisable to do so after the task has been completed in order to avoid priming student articulations and actions. Indicators introduced after task completion and prior to retrospective reflection might take the following forms.

Indicators of intelligence augmentation in screen recordings:

- Precision prompting in optimizing output
- Using AI as a successful control mechanism rather than a stand-alone authority
- Tool orchestration of gen AI with other resources to facilitate efficiency without sacrificing quality

Indicators of intelligence impairment in screen recordings:

- Single-shot acceptance of AI output without revision or cross-checking other resources
- Copying and pasting of content across interfaces without revision (when needed)
- Shallow initial prompts yielding problems or errors not addressed through follow-up prompting

Indicators of intelligence augmentation in think-aloud protocols:

- Articulations signaling direct engagement and critical reflection with gen AI in addressing problems
- Articulations predicting how gen AI might help prior to actual use
- Articulations involving rationales for why gen AI is assistive

Indicators of intelligence impairment in think-aloud protocols:

- Articulations reflecting superficial actions (“I’m entering a prompt”) rather than critical reflection
- Complete absence of articulation in conjunction with not detecting or addressing problems

- Articulations reflecting a loss of anchoring and detachment from the translation task and source content

Given the cognitively demanding nature of thinking out loud and subsequent manual retrospective analysis of on-screen phenomena and articulations, the translation task should be limited in text length and duration. Previous studies using this approach have used source content that is 100–150 words in length, yielding screen recordings that are approximately 20–30 minutes in duration.

In terms of concrete learning activities, obtained process protocols can be used for meeting various objectives. One activity could have students gauge their performance in conjunction with the aforementioned indicators to discern what augmentation and impairment look like in their particular case. A related learning activity could involve reverse engineering errors in the translation product in relation to what was articulated (or not articulated) and transpiring on-screen in conjunction with them. Students could be asked to reflect and report on the following: Was generative AI being used at the time the error occurred? If so, why did AI fall short and what might you have done differently to prevent the error? Do errors correlate with indicators of intelligence augmentation or intelligence impairment? What might have transpired differently had generative AI not been used in these instances?

References

- Appiah, Kwame Anthony. 2025. The age of de-skilling: Will AI stretch our minds—or stunt them? *The Atlantic*.
- Krüger, Ralph. 2025. Implementing generative artificial intelligence technologies in language industry workflows – A competence perspective. *Lebende Sprachen*, 70(1):11–38.
- Penet, JC, Joss Moorkens, and Masaru Yamada, Masaru, editors. 2026. *Teaching Translation in the Age of Generative AI: New Paradigm, New Learning?* Language Science Press, Berlin.
- Pym, Anthony, and Yu Hao. 2024. *How to Augment Language Skills: Generative AI and Machine Translation in Language Learning and Translator Training*. Routledge, London/New York.