

A Technical Curriculum on Language-Oriented Artificial Intelligence in Translation and Specialised Communication

Ralph Krüger

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

Abstract

This paper presents a technical curriculum on language-oriented artificial intelligence (AI) in the language and translation (L&T) industry. The curriculum aims to foster domain-specific technical AI literacy among stakeholders in the fields of translation and specialised communication by exposing them to the conceptual and technical/algorithmic foundations of modern language-oriented AI in an accessible way. The core curriculum focuses on 1) vector embeddings, 2) the technical foundations of neural networks, 3) tokenization and 4) transformer neural networks. It is intended to help users develop computational thinking as well as algorithmic awareness and algorithmic agency, ultimately contributing to their digital resilience in AI-driven work environments. The didactic suitability of the curriculum was tested in an AI-focused MA course at the Institute of Translation and Multilingual Communication at TH Köln. Results suggest the didactic effectiveness of the curriculum, but participant feedback indicates that it should be embedded into higher-level didactic scaffolding – e.g., in the form of lecturer support – in order to enable optimal learning conditions.

1 Introduction

The recent emergence of general-purpose AI (GPAI) technologies in the form of large language

models (LLMs) has widened the scope of automation in the L&T industry to a considerable extent, as discussed, e.g., in Lambson and Han (2025) and in Krüger and Angelone (forthcoming). For example, a single LLM can perform translation-technological tasks that were originally distributed over different expert technologies (e.g., intra- and interlingual machine translation, terminology extraction or quality estimation). Furthermore, LLMs can augment traditional translation technologies with new features that were previously beyond the scope of these technologies (e.g., augmenting traditional automatic translation quality control with meaning-based source- and target-segment comparisons) and they can (partially) automate tasks which largely resisted being automated to date (e.g., autonomous text production at high proficiency levels or intersemiotic machine translation/interpreting).

The current LLM-driven automation cycle in the L&T industry is shaping up to be both more extensive and more invasive than previous automation cycles, as evidenced, for example, by the *LangOps Manifesto*'s *AI-first* principle.¹ It could be claimed that – in order to retain high levels of agency and control in AI-first environments in the spirit of the human-centered AI approach² – L&T industry stakeholders require domain-specific AI literacy, as conceptualised, e.g., in the AI literacy framework proposed in Krüger (2024) and Krüger (forthcoming). In this framework, overarching AI literacy is further decomposed into the individual

¹“AI systems are already good enough for many use cases and their performance will continue to improve. We always start by trying AI solutions to solve language problems.” (<https://langops.org/our-manifesto/>)

²For a general overview of the concept of human-centered AI, see Shneiderman (2020), for a translation-specific perspective, see Jiménez-Crespo (2025).

subdimensions of technical, performance-focused, interaction-focused, implementation-focused and ethical/societal AI literacy, which are interrelated in various ways.

The curriculum proposed in this paper focuses on developing technical AI literacy, which involves knowledge of the basic operating principles of modern – currently to a large extent transformer-based – AI technologies. It is certainly true that, as noted by Bowker (2026, 73) with regard to translation didactics, “[i]t is not necessary for most translation educators or students to understand all the details of how an artificial neural network operates in order to use AI-based translation.” This view is complemented by Rivas Ginel and Moorkens (2025, 295), who found that translators’ (dis)trust in AI technologies such as LLMs is mostly calibrated by their use of these systems rather than by a deeper understanding of the systems’ technical foundations.

However, there are also several arguments in favour of fostering technical AI literacy in L&T industry stakeholders. Most importantly, technical AI literacy has an element of empowerment in that it helps in demystifying the black-box nature of modern machine-learning based AI technologies such as LLMs, a point which was previously made by Doherty and Kenny (2014, 296) in the context of statistical machine translation (SMT) and by Kenny (2019, 438) in the context of neural machine translation (NMT). AI-literate stakeholders familiar with the operating principles of these technologies may have a better-informed picture of the tasks these technologies can currently be used for and the level of quality at which these tasks can be completed (aspects related to performance- and implementation-focused AI literacy). This may also enable them to assume technology-oriented consulting roles in L&T industry contexts; see, e.g., the “consulting competence” in the post-editing competence model proposed by Nitzke et al. (2019, 248). At a higher level of abstraction, technical AI literacy may also help L&T stakeholders develop their computational thinking (Celik, 2023, 4) and, related to this, their algorithmic awareness (Gran et al., 2021) and algorithmic agency (Mills and Gutierrez, 2026, 42), ultimately contributing to their digital resilience (Kornacki and Pietrzak, 2024, 63, 65) in highly automated AI-first environments. These arguments in favour of technical AI literacy are supported, to some ex-

tent, by feedback on the didactic suitability of the technical curriculum as discussed in Section 3.3.

With its L&T stakeholder focus, the curriculum stands in a tradition of initiatives such as Multi-TraiNMT (Kenny, 2022), DataLit^{MT} (Krüger and Hackenbuchner, 2024), adaptMLLM (Lankford et al., 2023) and LT-LiDER (Moorkens et al., 2024).

2 The curriculum

The technical curriculum on language-oriented AI in translation and specialised communication is provided as a GitHub repository³, which consists of a series of Jupyter notebooks made available under a CC BY-SA 4.0 License (Attribution-ShareAlike 4.0 International) and hosted in Google’s Colaboratory (Colab) online environment. The core curriculum presented below is complete. Further notebooks addressing additional aspects of language-oriented AI relevant to the L&T industry will be added to the curriculum in the future.

2.1 Structure of the curriculum

The core curriculum consists of four sections, each assigned to a separate Jupyter notebook. These notebooks cover 1) vector embeddings, 2) the technical foundations of neural networks, 3) tokenization (with a focus on subword tokenization) and 4) the inner workings of transformer neural networks.

Vector embeddings – Representing linguistic data as numbers

The first notebook is concerned with vector embeddings and illustrates to users in a detailed way how linguistic data can be represented as numbers to be processed by neural networks. The notebook covers static/decontextualized and dynamic/contextualized word embeddings as well as sentence embeddings and multilingual embeddings. Users can load a pre-trained static word embedding model and train a small word embedding model themselves, explore individual word and sentence embeddings and embedding relations through visualisations and by calculating the Euclidean Distance and the Cosine Similarity between individual embeddings. The notebook also illustrates how transformer language models – taking BERT (Devlin et al., 2019) as an example – start with static/decontextualized embeddings in their initial embedding layer and convert these into

³https://github.com/ITMK/AI_Literacy/

dynamic/contextualized embeddings in their output layer (this process is explored in more detail in the fourth notebook concerned with transformer language models). This notebook is intended to meet L&T stakeholders where they are in their capacity as (aspiring) experts in natural language communication and guides them towards developing a more computational perspective on natural language that can be broadened and deepened in the further sections of the curriculum.

To rephrase this approach in didactic terms, the first notebook leads users “from the zone of current development to the zone of proximal development – the space where one cannot quite master a content/task of their own, but they can with the help of an expert” (Angelone, 2026, 33). The content covered in the subsequent notebooks is then mostly situated in this zone of proximal development. The expert guidance required in this zone can be embedded into the notebooks, to some extent, through adequate didactic scaffolding as discussed in Section 2.2. If the notebooks are used as learning resources in an instructor-led course, the lecturer can provide further expert guidance (see also the discussion in Section 3.2).

Technical foundations of neural networks – Building blocks, forward and backward pass

The second notebook addresses the technical foundations of neural networks and introduces users to the basic building blocks of these networks (neurons, trainable parameters, activation functions), their structure (input layer, hidden layers, output layer) and the flow of information through these networks (forward and backward pass). In this notebook, which intentionally favours didactic simplicity and accessibility over functional completeness, users can create a small experimental neural network in Python by defining the individual layers and the mathematical operations performed within and between these layers. Users can then simulate a simplified forward pass through this neural network, which illustrates how such a network may perform a translation of a short example sentence. The notebook also illustrates through a simplified backward pass how, in neural network training, the network’s parameters (weights and biases) would be updated via backpropagation in order to reduce network loss. Building upon the previous notebook, users learn how word vectors are combined into a matrix and how this matrix is processed by the in-

dividual layers of the neural network. The notebook is intended to further develop users’ computational thinking and algorithmic awareness by illustrating how natural language is actually processed and how natural language semantics is represented within a neural network. At the same time, it introduces users to foundational machine learning concepts and operations such as matrices, tensors, dot-product calculation, non-linear and softmax activation functions etc., which they will encounter again in the context of transformer neural networks.

Tokenization – Reducing the vocabulary size of neural networks

The third notebook is concerned with tokenization, which is employed in language models in order to reduce the size of the vocabulary to be learned. The notebook first guides users through word- and character-based tokenization and discusses the advantages and disadvantages of each method. From there, the notebook guides users through the process of subword-based tokenization, which combines the advantages of word- and character-based tokenization and is the tokenization method actually used in state-of-the-art language models. In this context, users are introduced to three major subword tokenization algorithms, i.e., Byte-Pair Encoding (BPE) (Gage, 1994), WordPiece (Schuster and Nakajima, 2012) and Unigram (Kudo, 2018). Users can then load a BPE, a WordPiece and a Unigram tokenizer, have them each tokenize an example sentence and explore how the three algorithms differ in how they tokenize words and which special characters they employ to designate word boundaries or intra-word splits. Then, users are introduced to the concept of token IDs, which are unique identifiers associated with each token in a language model. Armed with this knowledge, users can access the vocabulary of the language model GPT-2 (Radford et al., 2019) and explore which tokens are stored in this vocabulary. Finally, the notebook illustrates the steps involved in going from subword tokenization to vector embeddings to be processed by a language model.

Tokenization is covered after word embeddings and the technical foundations of neural networks since particularly subword tokenization introduces an additional level of complexity which may be confusing to users at first, but which should be less daunting once they are familiar with word embed-

dings and how these are processed by neural networks and once their computational thinking has evolved accordingly. Also, having encountered the encoder-only transformer model BERT in the first notebook and the decoder-only transformer model GPT-2⁴ in the current notebook, users are well-equipped to proceed to the next notebook, specifically concerned with transformer language models.

Transformer neural networks

The fourth notebook builds upon the foundations laid in the previous three notebooks and takes a deep dive into the inner workings of the transformer neural network architecture. Initial focus is placed on encoder-decoder vs. encoder-only vs. decoder-only models and on their respective areas of application and training regimes. In this context, the notebook illustrates how the Hugging Face transformers pipeline⁵ defaults to encoder-decoder models for sequence2sequence tasks such as text summarization or translation, to encoder-only models for natural language understanding tasks such as named entity recognition or question answering and to decoder-only models for natural language generation tasks such as text generation/completion. Then, the notebook explores in detail the individual components of the encoder and the decoder side of an encoder-decoder transformer language model. Here, the focus is placed on the original transformer architecture proposed by Vaswani et al. (2017) in the context of NMT research because, a) current LLMs are still strikingly similar to this original transformer architecture (Raschka, 2025a) and b) early transformer models are a bit simpler to understand than current LLMs “but still complex enough to give you a solid grasp of how modern transformer models work” (Raschka, 2025b). When more profound differences exist between the original transformer architecture and current LLMs (e.g., in terms of how positional embeddings are created, which activation functions are used or how layer normalization is performed), the notebook provides an ‘architecture update alert’, where it briefly discusses the cur-

rent techniques and points users to relevant literature. On the encoder side, the notebook first illustrates how inputs to the transformer are embedded (here, users can draw on their previous knowledge acquired in the notebook on vector embeddings) and how these embeddings are enriched with positional encoding vectors. Then, the notebook explores in detail the transformer self-attention process (splitting the original input embedding representation into query, key and value representations, dot-product calculation, scaling, padding/masking, softmax normalization etc.), which creates a dynamic/contextualized representation of the initially static/decontextualized input embedding representation. Armed with this conceptual and technical/algorithmic knowledge of transformer self-attention, users can then visualize and explore the self-attention process on their own by using the BertViz package (Vig, 2019). The discussion of the decoder side is more concise because the decoder is conceptually similar to the encoder, meaning that users will already be familiar with many of the components discussed here. Particular focus is placed on masked multi-head attention and on output generation in the final linear and softmax layers of the decoder. Here, users can have GPT-2 complete a text and visualize the top three token probabilities at each generation step. They can also access the user-friendly Hugging Face Decoding Visualizer⁶ to further explore greedy decoding and the Hugging Face Beam Search Visualizer⁷ to explore the more complex beam search decoding algorithm.

This core curriculum is currently being complemented by further notebooks covering training and adaptation strategies for language models as well as more recent developments in language-oriented AI, such as multimodal LLMs⁸, large reasoning and action models, small language models or knowledge-enhanced LLMs (e.g., via RAG or Knowledge Graphs). The notebooks of the core curriculum are also updated on an ongoing basis according to user feedback.

⁴The curriculum relies on older language models such as BERT and GPT-2 because they are functionally less complex than current state-of-the-art models and, more importantly, because they are considerably smaller and can thus be loaded into a Google Colab environment and used in this environment at reasonable speed.

⁵https://huggingface.co/docs/transformers/main_classes/pipelines

⁶https://huggingface.co/spaces/agents-course/decoding_visualizer

⁷https://huggingface.co/spaces/m-ric/beam_search_visualizer

⁸Including models capable of processing spoken language, which is intended to extend the applicability of the curriculum to the field of interpreting as well.

2.2 Didactic makeup of the Jupyter notebooks

The curriculum explores the didactic potential of Jupyter notebooks in introducing audiences with little to no programming knowledge and correspondingly low levels of computational thinking and algorithmic awareness to complex concepts and algorithmic processes from fields such as natural language processing, as investigated in a translation context in Krüger (2022). Jupyter notebooks support *literate computing* (Millman and Pérez, 2018) through a combination of (multi-modal) documentation sections providing conceptual information interleaved with executable code cells, which illustrate the algorithmic implementation of these concepts. This combination of documentation and executable code allows authors to use these notebooks to “tell an interactive, computational story” (Barba et al., 2019, 7) to their readers, with a didactic scaffolding tailored to specific target audiences. The target audience *L&T industry stakeholders* as referred to in this paper includes, e.g., students of MA programmes in translation, specialised communication (as central L&T industry fields) and related areas as well as professionals working in the L&T industry. It is assumed that these stakeholders have a professional interest in language-oriented AI but have not been exposed in detail to the technical foundations of these technologies. It is further assumed that these stakeholders have had moderate to no exposure to computer programming before, making them novices or advanced beginners in terms of computational thinking.

Figure 1 depicts a typical example of the didactic makeup of the Jupyter notebooks within the technical curriculum. The left screenshot is concerned with a forward pass of a simple example sentence through an experimental translation neural network and illustrates how the ‘semantics’ of a particular input word is processed and represented by a hidden layer in the network (excluding the activation function at this point). The right screenshot is concerned with self-attention in a transformer language model. The code cells implement BertViz to visualise the BERT model’s self-attention process performed on a predefined example sentence (which users are encouraged to change in order to experiment with their own sentences). The documentation sections explain how to interpret the BertViz visualisation. Short exer-

cises then invite users to engage more deeply with the topic, in the current example also highlighting points of contact with topics covered at a previous point in the curriculum.

3 Measuring the didactic suitability of the curriculum

The core curriculum presented above was tested in the MA course “Foundations of Artificial Intelligence for Specialised Communication” in the MA in Multilingual Specialised Communication and Specialised Translation programme (code: *MAFKÜ*) and the MA in Terminology and Language Technology programme (code: *MATS*) at TH Köln’s Institute of Translation and Multilingual Communication. The course was intended for first-year students of the two MA programmes and was also attended by three external guest students who are experienced professional translators working in public-sector language services (code: *professionals*). Participants completed a pre-test questionnaire (October 2025) and a post-test questionnaire (January 2026). The study was conducted in accordance with TH Köln’s Regulations for Safeguarding Good Scientific Practice⁹ and Code of Conduct for Research and Transfer¹⁰. The questionnaires were derived from the draft version of the TransAI Literacy Scale (TrAILS) proposed in Krüger (forthcoming) and were administered in German.¹¹

The pre-test questionnaire was completed by 24 participants. Of these, 12 (50%) were MAFKÜ students, 9 (37.5%) were MATS students and 3 (12.5%) were professionals (item 0.3). The latter group had a professional experience of more than 20 years, while the professional experience of the MATS and MAFKÜ students varied between <1 year (66.7%), 1 to <5 years (12.5%), 5 to <10 years (4.2%) and 10 to <20 years (4.2%) (item 0.2). The post-test questionnaire was completed by 15 participants, resulting in a survey attrition rate of 37.5%. The composition of the post-test group

⁹https://www.th-koeln.de/mam/bilder/forschung/2023-05-31_gwp_en.pdf

¹⁰https://www.th-koeln.de/mam/bilder/forschung/verhaltenskodex_maerz_2021_en-us.pdf

¹¹Draft version of TrAILS: <https://forms.gle/QdhPiP9D3F6rUca4A>
Pre-test questionnaire: <https://forms.gle/VNwL8LpNdsKh6HLd8>
Post-test questionnaire: <https://forms.gle/VBzvcFub3EtSLSA6>

Basically, this output represents how our input sentence / low neural networks was processed by the first neuron of the second hidden layer, after having been processed by all neurons of the first hidden layer. Since the sentence consists of four words, we also have four values in our vector. Let us zoom in again on the last word of this sentence (networks) and see how the first neuron of the second hidden layer processed it:

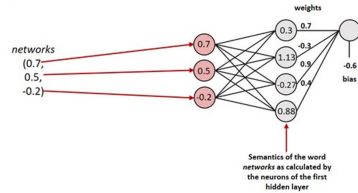


Fig. 13: How the first neuron of the second hidden layer processes the input word networks

Mathematically:

$$(0.3 \times 0.7) + (1.13 \times 0.9) + (-0.27 \times 0.4) + (-0.88 \times 0.4) + -0.6 = -0.62$$

In the figure above, the values inside the neurons of the first hidden layer are the values that the individual neurons calculated for the semantics of the word networks coming from the input layer (above, we saw that the first neuron of the first hidden layer calculated the value 0.3 for the semantics of the word networks). Here, this semantics flows from the first hidden layer via the weighted connections to the second hidden layer, thus being processed further. You see that the first neuron in the second hidden layer calculated the semantics of networks to be -0.62 , which corresponds to the last value of the vector above.

Time to exercise: How does the second neuron of the second hidden layer represent the semantics of the word neural? 🧐

```

inputs = tokenizer.encode("IHT doesn't feed on brainwave energy", return_tensors='pt')
outputs = model(inputs)
attention = outputs[-1]
tokens = tokenizer.convert_ids_to_tokens(inputs[0])

```

In BertViz, we can chose different views, of which we will explore two. We start with the *head view*, which lets us visualise attention in a specific attention head. Run the cell below to do so.

```

from bertviz import head_view
head_view(attention, tokens)

```

The BERT model has 12 attention layers and 12 attention heads per layer. Hence, in the window above, you can choose between 12 layers and for each layer activate up to 12 attention heads (the coloured squares) to be visualised. In the left column, our example sentence is split into query words (or query tokens, to be more precise). The right column represents the key words/tokens. If we place the cursor on one of the tokens in the query column, BertViz visualises the relevance of the key tokens for the contextualised representation of the query tokens (stronger lines indicate higher attention weights and vice versa).

Exercise 1 (a brief excursus): Can you identify which subword tokenisation method was used in our example above? 🧐 If you need to refresh your memory, take a look at our [notebook on tokenisation](#).

Exercise 2: Explore the self-attention distribution in individual layers/attention heads. Can you identify any interesting phenomena/patterns? 🧐

Figure 1: Example of the didactic makeup of the Jupyter notebooks – left: processing of embeddings in a hidden network layer (excluding activation); right: visualising and exploring self-attention in a transformer language model

shifted to 6 MATS students (40%), 5 MAFKÜ students (33.3%), 3 in-house professionals (20%, this group remained constant) and 1 freelance professional (6.7%).¹²

This shift is relevant for interpreting the results: The MATS curriculum includes a higher share of IT-related courses than the MAFKÜ curriculum (among other things, MATS students attended a course on markup languages, which ran parallel to the AI course and they will attend a separate course on Python programming in the second semester). Since attrition was considerably higher among MAFKÜ students compared to MATS students (58.33% vs. 33.33%), the post-test group likely exhibits a higher baseline of IT competence and technical disposition than the initial pre-test group. This is supported by the participants' self-assessment related to IT competence (item 0.6) and interest in the technical foundations of language-oriented AI (item 0.5). The share of participants reporting high or very high IT competence almost doubled from 20.9% in the pre-test to 40% in the post-test. Similarly, the share of participants reporting a high or very high interest in the technical foundations of language-oriented AI increased from 70.8% in the pre-test to 80% in the post-

¹²Since the study did not use constant participant IDs for pre- and post-test, it is unclear whether the participant who identified as a freelance professional in the post-test questionnaire identified as a MAFKÜ or a MATS student in the pre-test questionnaire (both questionnaires asked participants for their 'primary role' [item 0.3]; no further external guest students joined the course between the pre-test and the post-test). The attrition rate may stem from the course being offered every two semesters; consequently, some enrolled students may have decided to defer their participation to a future term.

test.¹³

3.1 Overall didactic effectiveness of the curriculum

The participants' language-oriented AI knowledge was measured using three questionnaire items (items 1.1 to 1.3) with an 11-point Likert scale ranging from 0 = *Hardly/Not at all* to 10 = *Perfectly* without intermediate scale labelling.¹⁴ The questions were asked in both the pre- and the post-test in order to measure participants' self-reported knowledge gain. Results are presented in table 1.

We focus here on the composite index, which represents the mean of the responses to items 1.1 to 1.3 for each participant (serving as an aggregate of participants' self-assessed knowledge of the technical foundations of language-oriented AI). The increase in participants' self-assessed technical knowledge from $M = 3.72$ ($SD = 2.13$, $N = 24$) in the pre-test to $M = 6.76$ ($SD = 1.43$, $N = 15$) in the post-test is highly statistically significant ($p < 0.001$)¹⁵, with a very large effect size

¹³There may be a combination of a selection effect (higher retention rate of students with a higher initial IT affinity; [very] high IT competence pre-test MAFKÜ group: 8.33% vs pre-test MATS group: 22.22% vs. pre-test professionals: 66.67%) and a treatment effect (increase in IT competence through contents covered in the technical curriculum) at work here. This will be discussed in more detail in Section 3.1.

¹⁴*I can explain the basics of how modern language-oriented AI technologies work* (item 1.1), *I can explain how modern language-oriented AI technologies are trained and fine-tuned* (item 1.2), *I know if language-oriented AI supports the translation technologies I use and, if so, how* (item 1.3).

¹⁵Welch's t-test for independent samples was used to account for the lack of consistent participant IDs across pre- and post-test and to address unequal group sizes and variances

Knowledge area	Pre-test ($M \pm SD$)	Post-test ($M \pm SD$)	p -value	Cohen's d
1.1 Basic operating principle	3.67 ± 2.24	6.73 ± 1.33	< 0.001	1.58
1.2 Training/Fine-tuning	3.04 ± 2.18	5.87 ± 1.77	< 0.001	1.39
1.3 AI-supported trans tech	4.46 ± 2.77	7.67 ± 2.09	< 0.001	1.27
Composite index	3.72 ± 2.13	6.76 ± 1.43	< 0.001	1.60

$N_{Pre} = 24, N_{Post} = 15$, 11-point Likert scale from 0 (*Hardly/Not at all*) to 10 (*Perfectly*)

Table 1: Self-reported knowledge gain of participants with regard to the technical foundations of language-oriented AI

of $d = 1.60$. This suggests a high overall didactic effectiveness of the technical curriculum, taking into account that generalizability may be limited due to the relatively small convenience sample size.

The post-test questionnaire also included a retrospective question where participants were asked to reassess their technical knowledge of language-oriented AI at the beginning of the course.¹⁶ This was done to account for a potential Dunning-Kruger effect at the onset of the study. The results of this retrospective self-assessment ($M = 2.93$, $SD = 2.19$) were notably lower than the initial self-assessment ($M = 3.72$, $SD = 2.13$), suggesting a response shift. As the curriculum introduced participants in detail to the technical foundations of language-oriented AI, they likely became aware of previous knowledge gaps that may have skewed their initial self-assessment. This indicates that the knowledge gain associated with the technical curriculum was even more pronounced than the initial pre-post comparison suggests. The response shift may also indicate an increase in participants' metacognitive awareness in terms of a more realistic picture of their own level of competence with regard to language-oriented AI.¹⁷

As mentioned above, the interpretation of results must take into account a potential selection effect (students with higher initial IT affinity completing the course and students with lower initial IT affinity leaving the course). However, both the observed response shift and the very large effect size of $d = 1.60$ (or $d = 2.07$ for the retrospective self-assessment of $M = 2.93$, $SD = 2.19$)

(*scipy.stats.ttest_ind, equal_var=False*).

¹⁶Looking back at the beginning of the course, how would you now assess your knowledge of the technical foundations of language-oriented AI at that time? (item 1.4, 11-point likert scale ranging from 0 = *Hardly any/No knowledge* to 10 = *Complete/Perfect knowledge*)

¹⁷Metacognition has also been discussed as a dimension of AI literacy acquisition in that it “[enables] individuals to reflect on their learning processes and understand how AI can enhance their cognitive abilities.” (Tadimalla and Maher, 2024, 3)

provide evidence for a genuine treatment effect. These factors indicate that the self-reported knowledge gain indeed resulted from participants' engagement with the technical topics covered in the curriculum, rather than being merely an effect of participant self-selection.

3.2 Didactic potential of Jupyter notebooks

The post-test questionnaire also asked participants to assess the didactic potential of Jupyter notebooks in the context of the technical curriculum. When participants were asked whether these notebooks were a suitable didactic instrument for covering the technical course contents (item 2.1), 12 (80%) strongly agreed, 1 (6.7%) agreed and 2 (13.3%) neither agreed nor disagreed. This largely confirms the results reported in Krüger (2022, 519–520) concerning the didactic potential of Jupyter notebooks in translation technology teaching. When asked specifically whether the combination of multimodal documentation (text + figures + linked videos) and executable code cells helped participants to understand the technical contents of the curriculum (item 2.2), 11 (73.3%) strongly agreed, 3 (20%) agreed and 1 (6.7%) neither agreed nor disagreed, which can be taken as further evidence in favour of the notebooks' didactic suitability.

Participants in the post-test study could also provide open comments on the didactic potential of the Jupyter notebooks for covering the technical course contents (five participants provided such comments). Overall, feedback was positive. E.g., according to a MATS student, the notebooks *illustrated the content very clearly and generally helped me to review the lecture material at my leisure (the notebooks were also very detailed, which was also helpful)*.¹⁸ A MAFKÜ student found *the structure and detail of the descriptions/explanations very helpful. Better than slides*

¹⁸Participants provided all open comments in German. The comments were machine-translated into English using DeepL and post-edited by the author.

with bullet points, because it gives you the opportunity to read up on something that has already been explained in detail by the lecturer in the seminar. The same student reported having difficulties saving the notebooks locally (which could have been achieved through adequate lecturer support). Another MAFKÜ student made an interesting point regarding further didactic scaffolding in addition to the notebooks: *In the context of our course and your explanations of the content of the lines of code, I found these very helpful. However, I'm not so sure whether I—or other users who have no prior knowledge of coding—would have been able to learn well with these alone, especially since the relevant information was often hidden in long lines of code.* This indicates that, in the zone of proximal development, expert guidance cannot be fully embedded into the notebooks themselves and should also be provided by the lecturer (in this specific case by highlighting relevant portions of longer and potentially complex code cells). Future iterations of the AI course will also introduce students to using LLMs as coding assistants for producing, highlighting, explaining, changing or debugging computer code, as discussed in a translation context by Krüger et al. (forthcoming). Such LLM integration may offer further didactic potential, e.g., in terms of using these models as general learning assistants facilitating deliberate practice (Angelone, 2026) in addition to lecturer support or in self-learning scenarios where no such lecturer support is available.¹⁹

3.3 Further feedback on the technical curriculum

When asked whether the progression of the course contents from 1) vector embeddings via 2) the technical foundations of neural networks and 3) subword tokenization to 4) transformer neural networks was suitable from a didactic point of view (post-test questionnaire, item 3.1), 12 participants (80%) found the order of topics highly suitable and 3 (20%) found it rather suitable. In an open comment, one of the professional participants noted that they would have liked a final overview of how the different AI concepts covered in the course re-

late to each other. This will be implemented in a future version of the curriculum.

Finally, participants were asked to provide concluding feedback (in the form of open comments) on the AI course and the technical curriculum, e.g., in terms of the usefulness of the course in the context of their MA programme or with a view to participants' (future) careers, in terms of the appropriateness of the technical level at which the content was covered or in terms of any content that was missing or was covered too briefly or superficially (item 4.1). Seven participants provided concluding feedback.

The technical level was considered demanding (particularly in the context of transformer language models), but participants largely agreed that it could be mastered using the didactic resources provided. A MAFKÜ participant noted: *Although I consider the course to be challenging, I find that the content is explained very clearly and from a perspective that is easy for me to understand (as a translator).* A professional translator participant commented that, at times, there was a risk of getting lost in too much technical detail. Another MAFKÜ participant responded that, although the lecturer always encouraged students to ask questions if they did not understand something, it was sometimes difficult to put such questions into words in the first place. These issues will have to be addressed in future iterations of the course, also taking into account the didactic potential of LLMs as coding or general learning assistants as discussed above.

Overall, participants acknowledged the relevance of the course and the curriculum for their (later) professional career. A MAFKÜ student noted that even if a technical understanding of language-oriented AI was not directly relevant in students' future careers, *I believe that this background knowledge gives participants a certain degree of confidence.* A professional translator noted that security restrictions currently prevented them from using language-oriented AI technologies extensively at their workplace but that *I can at least discuss with our IT department what is important for language service providers, make suggestions and discuss things on an equal footing.* These comments support the hypothesis stated in Section 1 that technical AI literacy as acquired through the technical curriculum includes an empowerment dimension and may strengthen L&T stakeholders'

¹⁹LLM's potential as learning assistants may be further reinforced by the possibility to prompt these models to tailor the "level of explanatory ambition" (Engberg, 2025, 82) to individual knowledge requirements (e.g., explaining the concept of a backward pass through a neural network to MSc students in information technology vs. MA students in translation/specialised communication).

(algorithmic) agency (e.g., by discussing AI ‘on an equal footing’ with IT experts and/or consulting on workflow design) and may contribute to their overall digital resilience (in terms of confidence) in highly automated environments.

4 Conclusion and outlook

This paper presented a technical curriculum on language-oriented AI in translation and specialised communication, which stands in the tradition of other L&T stakeholder-specific initiatives such as MultiTraiNMT, DataLit^{MT}, adaptMLLM and LT-LiDER. The overall didactic suitability of the curriculum was tentatively confirmed in an AI-focused MA course at TH Köln. Participants acknowledged the relevance of the course and the technical curriculum for their (later) professional careers, with individual participant feedback supporting the hypothesis that technical AI literacy may strengthen L&T stakeholders’ (algorithmic) agency and digital resilience in professional high-automation contexts. Participants also confirmed the didactic potential of Jupyter notebooks for covering the technical course contents. However, participant feedback also indicated that the technical content could be overwhelming at times or was only manageable through additional lecturer support. This indicates that, in the zone of proximal development, where learners are challenged by the novelty and/or complexity of content, the required expert guidance cannot be fully embedded into the Jupyter notebooks and should also be provided through higher-level didactic scaffolding, ideally in terms of lecturer support. Participant feedback will be implemented into the curriculum on an ongoing basis in order to tailor it further to the specific knowledge requirements of the target audience. This is intended to strengthen the curriculum’s *human-centered explainable AI* (HCXAI) dimension, which aims to place non-expert users at the center of AI explainability (Ridley, 2025, 98). Also, future work will investigate strategies for introducing LLMs as coding or general learning assistants into the curriculum, which could facilitate deliberate practice in self-learning scenarios.

Acknowledgements: The author would like to thank his colleagues Janiça Hackenbuchner and Erik Angelone for their valuable comments and suggestions on the manuscript version of this paper.

References

- Angelone, Erik. 2026. Generative AI as a facilitator of deliberate practice in translator training. In Penet, JC, Joss Moorkens, and Masaru Yamada, editors, *Teaching Translation in the Age of Generative AI: New Paradigm, New Learning?*, page 27–47. Language Science Press, Berlin.
- Barba, Lorena A., Lecia J. Barker, Douglas S. Blank, Jed Brown, Allen B. Downey, Timothy George, Lindsey J. Heagy, Kyle T. Mandli, Jason K. Moore, David Lippert, Kyle E. Niemeyer, Ryan R. Watkins, Richard H. West, Elizabeth Wickes, Carol Willing, and Michael Zingale. 2019. *Teaching and Learning with Jupyter*.
- Bowker, Lynne. 2026. Teaching translation students about data in the age of generative AI. In Penet, JC, Joss Moorkens, and Masaru Yamada, editors, *Teaching Translation in the Age of Generative AI: New Paradigm, New Learning?*, pages 67–85. Language Science Press, Berlin.
- Celik, Ismail. 2023. Exploring the determinants of artificial intelligence (AI) literacy: Digital divide, computational thinking, cognitive absorption. *Telematics and Informatics*, 83:1–11.
- Devlin, Jacob, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. 2019. BERT: Pre-training of deep bidirectional transformers for language understanding. In Burstein, Jill, Christy Doran, and Tamar Solorio, editors, *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers)*, pages 4171–4186, Minneapolis. Association for Computational Linguistics.
- Doherty, Stephen and Dorothy Kenny. 2014. The design and evaluation of a statistical machine translation syllabus for translation students. *The Interpreter and Translator Trainer*, 8(2):295–314.
- Engberg, Jan. 2025. Specialised communication and cognition. In Roelcke, Thorsten, Ruth Breeze, and Jan Engberg, editors, *Specialized Communication. An International Handbook*, page 67–86. De Gruyter Mouton, Berlin/Boston.
- Gage, Philip. 1994. A new algorithm for data compression. *The C Users Journal*, 12(2):23–38.
- Gran, Anne-Britt, Peter Booth, and Taina Bucher. 2021. To be or not to be algorithm aware: A question of a new digital divide? *Information, Communication & Society*, 24(12):1779–1796.
- Jiménez-Crespo, Miguel A. 2025. Human-centeredness in translation. Advancing translation studies in a human-centered AI era. *InContext*, 5(1):7–17.
- Kenny, Dorothy. 2019. Machine translation. In Rawling, Piers and Philip Wilson, editors, *The Routledge*

- Handbook of Translation and Philosophy*, pages 428–445. Routledge, London/New York.
- Kenny, Dorothy, editor. 2022. *Machine Translation for Everyone*. Language Science Press, Berlin.
- Kornacki, Michał and Paulina Pietrzak. 2024. *Hybrid Workflows in Translation: Integrating GenAI into Translator Training*. Routledge, London/New York.
- Krüger, Ralph and Janiça Hackenbuchner. 2024. A competence matrix for machine translation-oriented data literacy teaching. *Target*, 36(2):245–275.
- Krüger, Ralph and Erik Angelone. forthcoming. The AI age in the language and translation industry – Opportunities for automation and literacy requirements. In Jiménez-Crespo, Miguel A. and Vanessa Enríquez Raído, editors, *The Routledge Handbook of Translation and Technology (2nd ed.)*. Routledge, London/New York.
- Krüger, Ralph, Sergi Álvarez Vidal, and Janiça Hackenbuchner. forthcoming. Bridging the digital divide – Making Python and the Python ecosystem accessible to translators. In Castilho, Sheila, Dorothy Kenny, Joss Moorkens, and María Isabel Rivas Ginel, editors, *Developing AI Literacy for Translation*. Language Science Press, Berlin.
- Krüger, Ralph. 2022. Using Jupyter notebooks as didactic instruments in translation technology teaching. *The Interpreter and Translator Trainer*, 16(4):503–523.
- Krüger, Ralph. 2024. Outline of an artificial intelligence literacy framework for translation, interpreting and specialised communication. *Lublin Studies in Modern Languages and Literature*, 48(3):11–23.
- Krüger, Ralph. forthcoming. Artificial intelligence literacy for the language and translation industry – Conceptual foundations, operationalisation, acquisition, measurement. In Castilho, Sheila, Dorothy Kenny, Joss Moorkens, and María Isabel Rivas Ginel, editors, *Developing AI Literacy for Translation*. Language Science Press, Berlin.
- Kudo, Taku. 2018. Subword regularization: Improving neural network translation models with multiple subword candidates. In Gurevych, Iryna and Yusuke Miyao, editors, *Proceedings of the 56th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 66–75, Melbourne. Association for Computational Linguistics.
- Lambson, Nick and Lintao Han. 2025. Localization in the AI era. In Sun, Sanjun, Kanglong Liu, and Riccardo Moratto, editors, *Translation studies in the age of artificial intelligence*, pages 62–84. Routledge, London/New York.
- Lankford, Séamus, Haithem Afli, and Andy Way. 2023. adaptMLLM: Fine-tuning multilingual language models on low-resource languages with integrated LLM playgrounds. *Information*, 14(12):1–24.
- Millman, K. Jarrod and Fernando Pérez. 2018. Developing open-source scientific practice. In Stodden, Victoria, Friedrich Leisch, and Roger D. Peng, editors, *Implementing Reproducible Research*, pages 149–184. CRC Press, Boca Raton.
- Mills, Kathy A. and Amanda Gutierrez. 2026. *Critical Literacy in an AI World*. Routledge, London/New York.
- Moorkens, Joss, Pilar Sánchez-Gijón, Esther Simon, Mireia Urpí, Nora Aranberri, Dragoş Ciobanu, Ana Guerberof-Arenas, Janiça Hackenbuchner, Dorothy Kenny, Ralph Krüger, Miguel Rios, Isabel Ginel, Caroline Rossi, Alina Secară, and Antonio Toral. 2024. Literacy in digital environments and resources (LT-LiDER). In Scarton, Carolina, Charlotte Prescott, Chris Bayliss, Chris Oakley, Joanna Wright, Stuart Wrigley, Xingyi Song, Edward Gow-Smith, Mikel Forcada, and Helena Moniz, editors, *Proceedings of the 25th Annual Conference of the European Association for Machine Translation (Volume 2)*, pages 55–56, Sheffield, UK. European Association for Machine Translation (EAMT).
- Nitzke, Jean, Silvia Hansen-Schirra, and Carmen Canfora. 2019. Risk management and post-editing competence. *Journal of Specialised Translation*, 31:239–259.
- Radford, Alec, Jeff Wu, Rewon Child, David Luan, Dario Amodei, and Ilya Sutskever. 2019. Language models are unsupervised multitask learners. *OpenAI*.
- Raschka, Sebastian. 2025a. The big LLM architecture comparison. From DeepSeek-V3 to Kimi K2: A look at modern LLM architecture design. *Ahead of AI Blog*.
- Raschka, Sebastian. 2025b. From GPT-2 to gpt-oss: Analyzing the architectural advances and how they stack up against Qwen3. *Ahead of AI Blog*.
- Ridley, Michael. 2025. Human-centered explainable artificial intelligence: An annual review of information science and technology (ARIST) paper. *Journal of the Association for Information Science and Technology*, 76(1):98–120.
- Rivas Ginel, María Isabel and Joss Moorkens. 2025. Translators’ trust and distrust in times of GenAI. *Translation Studies*, 18(2):283–299.
- Schuster, Mike and Kaisuke Nakajima. 2012. Japanese and korean voice search. In *2012 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, pages 5149–5152, Kyoto.

- Shneiderman, Ben. 2020. Human-centered artificial intelligence: Three fresh ideas. *AIS Transactions on Human-Computer Interaction*, 12(3):109–124.
- Tadimalla, Sri Yash and Mary Lou Maher. 2024. AI literacy for all: Adjustable interdisciplinary socio-technical curriculum. In Phillips, Margaret, Jason B. Reed, and Dave Zwicky, editors, *2024 IEEE Frontiers in Education Conference (FIE)*, pages 1–9, Los Alamitos. IEEE Computer Society.
- Vaswani, Ashish, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N Gomez, Łukasz Kaiser, and Illia Polosukhin. 2017. Attention is all you need. In Guyon, I., U. Von Luxburg, S. Bengio, H. Wallach, R. Fergus, S. Vishwanathan, and R. Garnett, editors, *Advances in Neural Information Processing Systems*, volume 30, pages 1–11. Curran Associates, Inc.
- Vig, Jesse. 2019. A multiscale visualization of attention in the transformer model. In Costa-jussà, M. R. and E. Alfonseca, editors, *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics: System Demonstrations*, pages 37–42, Florence. Association for Computational Linguistics.