

ORIGINAL ARTICLE OPEN ACCESS

Leveraging Large Language Models to Generate Course-Specific Semantically Annotated Learning Objects

Dominic Lohr¹  | Marc Berges¹  | Abhishek Chugh²  | Michael Kohlhasse³  | Dennis Müller³ 

¹Professorship for Computer Science Education, Erlangen, Germany | ²Sophize.Org, Bengaluru, India | ³Professorship for Knowledge Representation and Management, Erlangen, Germany

Correspondence: Dominic Lohr (dominic.lohr@fau.de)

Received: 25 March 2024 | **Revised:** 19 August 2024 | **Accepted:** 24 November 2024

Funding: This work was supported by Federal Ministry of Education and Research Germany (Grant: 16DHBKI089).

Keywords: automatic question generation | computer science | generative AI | GPT-4 | large language models | retrieval-augmented generation

ABSTRACT

Background: Over the past few decades, the process and methodology of automatic question generation (AQG) have undergone significant transformations. Recent progress in generative natural language models has opened up new potential in the generation of educational content.

Objectives: This paper explores the potential of large language models (LLMs) for generating computer science questions that are sufficiently annotated for automatic learner model updates, are fully situated in the context of a particular course and address the cognitive dimension *understand*.

Methods: Unlike previous attempts that might use basic methods such as ChatGPT, our approach involves more targeted strategies such as retrieval-augmented generation (RAG) to produce contextually relevant and pedagogically meaningful learning objects.

Results and Conclusions: Our results show that generating structural, semantic annotations works well. However, this success was not reflected in the case of relational annotations. The quality of the generated questions often did not meet educational standards, highlighting that although LLMs can contribute to the pool of learning materials, their current level of performance requires significant human intervention to refine and validate the generated content.

1 | Introduction

The determinants of learning success have been extensively researched in various disciplines (Hattie 2009; Kirby and Lawson 2012). Numerous empirical studies support the hypothesis that tailoring learning materials to learner's needs significantly increases the effectiveness of learning outcomes (Sweller 1988; Sweller, van Merriënboer, and Paas 1998). The mastery learning theory by Bloom (1968) or the personalised system of instruction theory by Keller (1968) postulate teaching methodologies that highly build on individualised learning materials and assessment on a very fine-grained level. However,

resource limitations and an increasingly diverse educational landscape challenge teachers and content creators in general.

Adaptive learning systems such as ALEA (Kruse et al. 2023) promise to address this need by delivering Learning Objects (LOs), such as definitions, examples or questions tailored to individual learners' specific prior knowledge, competencies and preferences. One way of realising this is by using semantically annotated LOs that allow for determining their prerequisites, which concepts they address and which competencies they intend to foster. In conjunction with a model of a learner's current knowledge and competencies, this allows for many

This is an open access article under the terms of the [Creative Commons Attribution-NonCommercial](https://creativecommons.org/licenses/by-nc/4.0/) License, which permits use, distribution and reproduction in any medium, provided the original work is properly cited and is not used for commercial purposes.

© 2024 The Author(s). *Journal of Computer Assisted Learning* published by John Wiley & Sons Ltd.

Summary

- What is already known about this topic
 - While automatic question generation (AQG) is extensively studied, it mostly produces questions assessing factual recall, targeting the cognitive dimension *remember*.
 - Generation of questions that align with specific course conventions and context remains a significant challenge.
 - Semantically annotated learning objects are crucial for usage in user-adaptive learning assistants, but it is a time-consuming task to create them by hand.
- What this paper adds
 - Exploring the potential of large language models for generating computer science questions that are sufficiently annotated for automatic learner model updates, fully situated in the context of particular courses and address the cognitive dimension *understand*.
- Implications for practice and/or policy
 - As generated questions often contain (minor) errors, they are not ‘ready to use’, but the effort required to correct them is usually less than creating them completely by hand.
 - The automatically generated questions can serve as inspiration for educators and contribute to a more heterogeneous pool of learning objects.

valuable services, such as the automatic generation of *flashcards* or so-called *guided tours*: individually selected sequences of LOs (‘learning paths’) that can be automatically created by learners on demand by a simple click, to open up a specific concept in a specific cognitive dimension. To make this possible, the pool of semantically annotated LOs must be just as diverse as the educational biographies of the learners. However, developing sufficiently annotated LOs requires huge amounts of manual labour and expertise. This bottleneck underlines the need for innovative solutions that can generate customised learning materials automatically, thus becoming a focus of current research in computer science education. Attempts to (partially) automate the creation of quiz questions can be found in the literature as early as the 1970s (Wolfe 1976).

Studies on the role of artificial intelligence in education point out that a significant hurdle is the deficiency of suitable learning materials for individualised and adaptable learning (Denny et al. 2023), a problem that may be mitigated by the capabilities of large language models (LLMs). Research in computing education is exploring the capabilities of these models to generate educational content that is both contextually appropriate and educationally demanding, offering a promising solution to the constraints associated with manually creating content. However, the effectiveness of LLM-generated educational content is often limited by their lack of focus on specific course content and the learner’s current understanding. The generated LOs often do not fit within the intended educational framework or meet the diverse needs of individual learners.

This paper presents the results of experiments on generating semantically annotated quiz questions using a state-of-the-art LLM. In particular, we investigate the question of to what extent LLMs can be used to generate questions in the domain of university-level computer science (CS) that are didactically valuable, are sufficiently annotated to allow for the above selection process, can ideally be graded by a software system (e.g., multiple-choice or fill-in-the-blanks questions) to automatically and immediately increase the accuracy of the system’s learner models and are entirely situated in the context of a particular university-level course concerning terminology and notations used. Unlike earlier similar studies (Song et al. 2024; Sarsa et al. 2022), our task requires extensive context, rendering it unsuitable for naive approaches, for example, using ChatGPT. Instead, we use more targeted techniques beyond ‘prompt engineering’, such as retrieval-augmented generation (RAG).

Overview Section 4 details our experimental methodology, including the considerations behind our choice of semantic annotations (Section 4.1), the selection criteria for the LLM (Section 4.2), the design of our question generation pipeline (Section 4.3) and the evaluation framework used, including the criteria for assessing question quality and an expert-based survey methodology. The results of our experiments are presented and discussed in Section 5. Finally, Section 6 concludes the paper with a summary of our findings, implications for future research in automatic question generation (AQG) using LLMs, and potential pathways for enhancing the educational value of LLM-generated content in adaptive learning environments.

2 | Related Work

Over the past few decades, the process and methodology of AQG have undergone significant transformations, driven primarily by advancements in computational linguistics and artificial intelligence (AI). In the earlier stages, AQG relied heavily on rule-based systems that applied predefined templates and linguistic patterns to generate questions from text, requiring extensive manual crafting and domain-specific adjustments. First studies can be found by Wolfe (1976).

With the development of deep learning and neural network models, a transition to more sophisticated, context-aware systems can be recognised. These models, particularly sequence-to-sequence and transformer-based architectures, have enabled the generation of more nuanced, relevant and diverse questions by ‘understanding’ deeper semantic relationships within the text. Kurdi et al. conducted a systematic review of empirical research focused on addressing the issue of AQG within educational contexts (Kurdi et al. 2020). They thoroughly outlined various methods of generation, tasks and evaluation techniques found in studies between 2015 and early 2019. A standard method involves identifying sentences in the text sources with high information content using topic words or key phrases. The system then selects a keyword as the answer key, removes it to form a question stem and generates incorrect choices (distractors) using a clustering method without external data. These systems mainly produce questions testing *remember* *factual* knowledge, not *understanding*, a key focus of the work presented in this paper, and the results show that LLMs can potentially benefit

from semantics-based approaches to generate meaningful questions that are closely related to the source content.

Recent progress on generative natural language models has opened up new potentials in the generation of educational content (Prather et al. 2023; Bhat et al. 2022; Wang et al. 2022; Wong, Park, and Looi 2024). In recent years, more and more approaches have been found to generate tasks in a single step using LLMs such as GPT-3 instead of dividing the AQG process into several sub-tasks. Yan et al. did a systematic scoping review of articles published since 2017 to pinpoint the current state of research on using LLMs (Yan et al. 2024). They identified content generation, including multiple-choice questions and feedback generation, as primary educational tasks that research aims to automate. McNichols et al. tested the effectiveness of LLM-based methods for automatic distractors and feedback generation using a real-world dataset of authentic student responses (McNichols et al. 2023). Their findings show that fine-tuning was ineffective and that other approaches than LLM-prompting need to be explored. Dijkstra et al. (2022) developed a quiz generator based on OpenAI's GPT-3 model, fine-tuned on text-quiz pairs to generate complete multiple-choice questions, with correct answers and distractors. They noted that while most of the generated questions were of acceptable quality, creating high-quality distractors was more challenging than generating question-and-answer pairs.

More specifically, LLM-based approaches have recently been applied in programming education (for an intensive literature review, see (Prather et al. 2023)). Sarsa et al. (2022) explored the capabilities of OpenAI's LLM *CodeX* to generate programming exercises and code explanations. Their results show that the generated questions and explanations were novel, sensible and, in some cases, ready to use. Tran et al. (2023) evaluated the capabilities of OpenAI's GPT-4 and GPT-3 models to generate isomorphic multiple-choice questions (MCQs) based on MCQ-stems from a question bank and an introductory computing course. Their findings underscore that the newer generation of LLMs outperforms older generative models in AQG.

2.1 | Research Gap

Reviewing the current landscape of research in AQG, several critical gaps emerge that necessitate further investigation. Current research shows that methods for generating closed-format tasks that target the cognitive dimension of *remember factual* knowledge work well. However, generating tasks that target deeper understanding, such as those that require comprehension or analytical skills, remains a significant challenge. We also could not find any research on generating questions containing semantic annotations. While it is already known that LLMs have the potential for generating feedback, and in particular code explanations, no work attempts to generate questions using LLMs that go beyond introductory courses in computing education. Kurdi et al. (2020) point out that there is little focus on controlling question parameters like difficulty and generating feedback. Since we want to incorporate the generated questions into an adaptive learning assistant, we have carefully considered which parameters are relevant for adaptive learning—such

as difficulty, cognitive dimension and prior knowledge—and want to explore how these can be integrated into the question generation process. We also did not find any studies that investigated generating questions using advanced LLM approaches such as RAG, as it is crucial for scalability beyond the original training data of the model.

3 | Requirements for the Generated Questions

We posit that AQG in the context of university courses poses additional challenges barely (or not at all) covered by the existing literature, especially in a domain such as *computer science*:

Firstly, we claim that in more abstract domains such as mathematics and math-related subfields of computer science, questions solely focusing on remembering facts, or rudimentary application exercises are less suitable for the outcomes of a university-level class. This makes designing appropriate questions a more challenging task because doing so requires a level of understanding of and experience with the learning material.

Secondly, although the topics covered in a particular course are common across different universities, the details are much less standardised (e.g., precise definitions, terminology and notational conventions). This entails the additional requirement that questions (whether automatically generated or not) need to be formulated and situated in line with the conventions in a particular course, which are also more frequently subject to change. This makes approaches based on large amounts of dedicated training data, which necessarily do not generalise beyond a specific course *and* instructor's preferences, unsuitable. Instead, we should be able to provide the generator with the relevant learning materials (course notes, slides, etc.) whenever we want to generate a new batch of questions.

Additionally, our goal is to utilise the generated questions in an adaptive learning assistant that is capable of automatically selecting suitable questions based on the associated *learning objective* and a particular student's (estimated) prior knowledge in the form of a *learner model*. Therefore, we need to annotate the questions with the relevant information for this selection process, namely:

1. The *concepts* the question is intended to test,
2. the *cognitive level* the question targets (modelled as levels in Bloom's revised taxonomy (Anderson and Krathwohl 2001)) and
3. the *prerequisite* concepts occurring in the question and the associated competencies that a user should have mastered for the question to be suitable.

While this is less difficult for an experienced user, providing these annotations is still time-consuming and potentially automatable. Therefore, we also investigate the extent to which a LLM can directly generate fully annotated questions.

Figure 1 shows an example of a fully annotated multiple-choice question created by hand in LATEX following the Y-model framework (Lohr et al. 2023). The relational semantic

```

1 \begin{sproblem}[title=Logic-Based Agent]
2   \usemodule[MiKoMH/AI]{logic/slides?wumpus-world}
3   \usemodule[smg lom/epistemology]{mod?inference}
4   \usemodule[MiKoMH/AI]{rational-agents/slides?learning-agent}
5   \usemodule[MiKoMH/AI]{logic/slides?logic-based-agent}
6   \usemodule[MiKoMH/AI]{csp/slides?backtracking-search}
7   \objective{remember}{logic-based agent}
8   \objective{understand}{logic-based agent}
9   %
10  Which of the following statements are true about the \sns{logic-based
11  agent} in the
12  \sn{Wumpus world}?
13  \begin{mcb}
14    \mcc{The \sn{agentenv?agent} represents the \sn{Wumpus world} as a
15    \sn{constraint network} and
16    ``thinks'' by \sn{backtracking search}.}
17    \mcc{T}The \sn{agentenv?agent} represents the \sn{Wumpus world} in
18    a world description language
19    (aka. logic) and through \sn{inference} finds out new \sn{
20    knowledge} about the world to base its
21    \sns{action} on.}
22    \mcc{The \sn{agentenv?agent} learns new facts about the world by
23    comparing the world
24    description to the external \sn{performance standard}.}
25  \end{mcb}
26 \end{sproblem}

```

FIGURE 1 | Source code and representation in the system.

annotations allow the learner to, for example, hover over concepts in the text (after conversion to HTML) and display definitions (this functionality can, of course, be deactivated if desired). A detailed description of the specific annotation schema can be found in Section 4.1.

Ensuring the quality of the generated tasks is a primary goal. As there is no standard metric in the literature for measuring the quality of questions (Ch and Saha 2020), we reviewed metric lists from literature reviews on AQG (Ch and Saha 2020; Mulla and Gharpure 2023) to identify relevant criteria that determine the quality of the questions we generate, namely:

- Correct technical language,
- appropriateness for a particular course context, including definitions, annotations and conventions,
- feasibility of solving the question with the provided details and available course materials,
- clarity and lack of ambiguity,
- relevance to achieving the intended learning outcomes,
- alignment with the specific format of the task and
- automatically gradable (closed format).

Additionally, feedback is one of the most influencing factors for learning success (Hattie and Timperley 2007), and recent experiments show the potential of LLMs for generating feedback in programming education (Lohr et al. 2024; Kiesler, Lohr, and Keuning 2023; Hellas et al. 2023; Balse et al. 2023; Matelsky et al. 2023). A further goal is to generate questions containing feedback that helps learners understand why a given answer is wrong.

The research questions are as follows:

Problem 2 of 8 (Logic-Based Agent)

Quiz Ended

1 pt

Which of the following statements are true about the logic-based agents in the Wumpus world?

☐ The agent represents the Wumpus world as a constraint network and "thinks" by backtracking search.

Wrong

Definition 0.1. A **constraint network** is a triple (V, D, C) , where V is a finite set of variables, $D := \{D_v | v \in V\}$ the set of their domains, and $C := \{C_{uv} \subseteq D_u \times D_v | u, v \in V \text{ and } u \neq v\}$ is a set of constraints with $C_{uv} = C_{vu}^{-1}$.

☒ The agent represents the Wumpus world as a constraint network and "thinks" by backtracking search. We call the undirected graph $(V, \{(u, v) \in V^2 | C_{uv} \neq \emptyset\})$, the **constraint graph** of γ . Inference finds out new knowledge about the world to base its actions on.

Correct

☐ The agent learns new facts about the world by comparing the world description to the external performance standard.

Wrong

Objectives

logic-based + agent

logic-based + agent

Preconditions

performance + standard, agent, action, knowledge, inference, Wumpus + world, backtracking + search, constraint + network, logic-based + agent

RQ1. To what extent can LLMs be used to generate university-suitable autogradable questions in CS education?

RQ2. To what extent can LLMs be used to annotate these questions semantically?

4 | Methodology

4.1 | Semantic Annotations in STEX

To semantically annotate learning objects—including questions—we use the STEX package (Kohlhase and Müller 2022) for LATEX. STEX uses an ontology based on OMDoc (Kohlhase 2006): Concepts are represented as *symbols* that can be introduced via the `\symdecl`-macro and can be related to each other in various ways. Symbols are always declared in *modules* (via the `smodule`-environment), which can *import* (the symbols of) other modules, adding another layer of relations and allowing for sharing concepts among large collections of disparate documents, thus enabling the collaborative and modular development of domains of knowledge as highly interrelated knowledge graphs independent of, and across, *presentational* context: symbols can have arbitrarily many definitions, formulations thereof and different notations, to accommodate author's preferences without duplicating the domain knowledge itself.

Symbols can be *referenced* in various ways; most importantly, the `\symref`-macro allows annotating arbitrary text as representing a particular symbol, and formal notations can be produced via dedicated *semantic macros* that additionally associate the notation with the corresponding symbol. Since the details are largely irrelevant to the topic of this paper, we refer to (Kohlhase and Müller n.d.) for details and Section 4.4 for examples.

More importantly, for our purposes, STEx also bundles the `problem`-package, providing dedicated markup macros and environments for various variants of quiz questions. Most importantly, it provides the `sproblem`-environment to annotate questions, within which we can use the `mcb` and `scb` environments for *multiple-* and *single-choice blocks*, respectively. As a third (autogradable) question type, the `\fillinsol` macro can be used to generate a blank box for *fill-in-the-blanks* questions. In all three cases, we can mark (or provide) the correct answer(s), add feedback text to be displayed if a particular answer (correct or wrong) is chosen and specify *grading actions* (i.e., set, add or deduct points) depending on the answers given. Additionally, we can specify the *learning objectives* and *prerequisites* as pairs of a *symbol* and one of the keywords `remember`, `understand`, `apply`, `analyse`, `evaluate` and `create`. Symbols referenced in the body of the question are automatically determined to be prerequisites with the cognitive dimension `remember`. For example, if a module declares the symbol `plus` for addition on integers, then the question ‘what is $2 + 2$?’ automatically has the prerequisite (`remember`, `plus`). We refer to such annotations as **relational annotations** and to those that do not relate text to some symbol (e.g., `fillinsol` or `sproblem`) as **structural annotations**.

Besides being LATEX compilable to pdf, we can use the $R_{\text{US}}\text{TEX}$ system (Müller 2023) to convert the STEx documents to HTML, preserving the semantic annotations in the form of attributes on the HTML nodes. Additionally, this assigns a globally unique URI to the document itself and every section, module, symbol and learning object therein, enabling searching, querying and referencing all of these afterward *across* documents.

Subsequently, we can embed various services that act on the semantic annotations directly into the HTML documents via JavaScript. Our learning assistant utilises this to render quiz questions as interactive components, evaluate answers provided by students, display the appropriate feedback provided and update a student’s learner model accordingly (see Figure 1).

While our usage of STEx is primarily motivated by our learning assistant, for the purposes of this paper, it offers additional advantages: Its underlying ontology, being designed around representing the semantics of mathematical statements formally, is consequently sufficiently expressive to subsume most semantic annotation systems, in particular those tailored specifically for quiz questions. Our results should, therefore, (all else being equal) generalise to other annotation schemes. Furthermore, since STEx inherits its syntax from LATEX, the LLM used does not need to be fine-tuned or explicitly prompted on the usage of a new or esoteric language—the huge amount of publicly available LATEX code implies that its syntax is well represented in the training data used for these models and can therefore be expected not to pose an issue for generation. Indeed, our experiments did not yield any LLM output with basic syntax errors.

4.2 | Model Selection

Our goals put notable constraints on the methodology for generating questions:

1. Targeting *understanding* rather than *remembering* factual knowledge implies that the model we use should be capable of ‘synthesising’ complex knowledge, excluding, for example, smaller language models less capable of ‘reasoning’.
2. Since the knowledge domain and preferred conventions should not be fixed, deliberately training a model on dedicated course materials is largely not feasible, which calls for *few-shot* approaches and
3. Since the latter point also entails providing possibly large amounts of course materials *online* during generation, the model used needs to be able to process large amounts of text at once. More precisely, it should allow for a large enough *context window*—the maximum number of tokens the model can consider in a single prompt/reply step.

OpenAI’s GPT-4 (OpenAI 2023) shows consistently better results across all tasks compared to alternative LLMs¹. In particular, we initially experimented with the free GPT-3.5 model (that powers the free version of ChatGPT), which quickly demonstrated very poor performance both concerning the questions generated and, more broadly, the ability to follow the instructions provided—meaning more prompt design was unlikely to improve the results significantly.

Similarly, most models (including almost all open-source models) primarily focus on short conversations and, therefore, only support a relatively small context window. For example, one of the currently most popular open-source language models, *LLaMA* (Touvron et al. 2023), can process a total input of 2048 tokens, whereas GPT-3.5 and GPT-4-Turbo have context windows of up to 16,385 and 128,000 tokens, respectively.²

We consequently opted for the (unfortunately commercial, closed source and proprietary) GTP-4-Turbo model as the one holding the most promise with respect to our criteria at the time of the experiments. Unfortunately, this also means we cannot provide access to a public instance of our pipeline since the API used to access the model is monitored and billed on a per-token basis.³

4.3 | Overview of the Generation Pipeline

To allow for the model to generate questions for a *specific course* regarding terminology, definitions and notations, we opted for a technique called RAG. This technique gives models access to additional information beyond their training data by querying an external knowledge base (such as a database or a web search engine) for results relevant to the specific prompt. It concatenates them to the prompt itself before passing it into the model.

To explain how we use RAG, we give a brief overview of the entire question generation pipeline. Note that we deliberately leave room for variation:

- A course instructor selects some concept (represented by an STEx symbol), a cognitive dimension, a document of course materials (i.e., the entire lecture notes) via its URI and additional parameters (e.g., number of questions, difficulty level, types of questions and a brief description of the course topics).

- The system replaces placeholders in a generic ‘*master prompt*’ (see Section 4.4) by the parameters provided.
- The system then selects those text fragments in the provided course materials that directly relate to the chosen concept—for example, definitions, examples and surrounding remarks—on the basis of the semantic annotations in the document and appends them to the master prompt. Since symbols are intrinsically relational, we can, in principle, also add materials for *dependent concepts* down to some cut-off point. In practice, we currently pick the entire chapter (or section or subsection) in which the concept is introduced.
- The final prompt is passed to the LLM and its output is presented to the instructor.

We focus here on the scenario where all parameters—in particular, the concept, cognitive dimension and number of questions—are explicitly fixed in the prompt, as this allows for better evaluation of the results of our experiments. However, we note that we can easily generalise to broader applications by automatically selecting some of the parameters. For example, we can generate questions for a whole *chapter* along *all* cognitive dimensions by determining the concepts introduced therein (which we can do automatically based on semantic annotations) and repeatedly running the above pipeline with, for example, randomly varied parameters.

We also chose not to split the task into multiple prompts for specific purposes. For example, we could conceptually use three distinct prompts (or even distinct *models*) for 1. generating questions, 2. adding feedback for students and 3. introducing semantic annotations. This split would be natural in a non-LLM-based approach since all three steps require distinct methodologies and tools. However, when using the same model for all three steps, it is less clear whether splitting the task is advantageous. We leave the question of whether it is for future work but note that our initial (and admittedly superficial) experiments in that direction did not yield noticeably different results. Furthermore, experiments with instructing the model to (paraphrased) *first* ‘think of a good question, summarize it, and *then* state the question as fully annotated STEX’—conceptually splitting the task into two separate steps—did not seem to make a clear difference in the quality of the output either. This is to some extent unexpected since it closely corresponds to the popular ‘think step by step’-instruction that is considered to generally improve results in LLM prompts related to reasoning.

4.4 | Prompt Design

Despite ongoing attempts to declare ‘prompt engineering’ a marketable skill, few hard principles for designing LLM prompts can be consistently well supported empirically. The most important rules of thumb can be briefly summarised thusly (Zamfirescu-Pereira et al. 2023; OpenAI 2024):

1. Be *specific* in what the output should look like,

2. be *detailed* by including any and all information and requirements relevant to the expected output,
3. provide *examples* in the prompt, effectively transforming the task from a *zero-shot* to a *few-shot* approach, and, most importantly,
4. *iterate* by repeatedly testing, observing problems with the outputs and modifying the prompt to discourage these problems in subsequent iterations.

Paraphrasing Zamfirescu-Pereira et al. (2023), it helps thinking of the language model as an *interpreter* and prompting as analogous to *programming* in the sense that it requires clear and detailed instructions to an entity keen to ‘do what you say, not what you (clearly) *mean*’.

Consequently, our final prompt is in many places the result of flaws occurring in previous iterations, which we will note where appropriate.

Our prompt starts with a brief summary of what we expect the model to do:

```
Your task is to generate quiz questions that
evaluate the competency of university stu-
dents with respect to a number of particular
concepts in a particular university course.
You will be given
- the names of the concepts,
- the name of the course,
- a cognitive dimension to test for,
- ...
```

We then clarify what we mean by *cognitive dimension*. Note that the huge amounts of data the model was trained on also contain many texts about concepts in education; particularly Bloom's revised taxonomy (Anderson and Krathwohl 2001). Nevertheless, we prefer to be precise:

```
We use the term "cognitive dimension" ac-
cording to Bloom's taxonomy. It is provided
as one of the strings "remember", "under-
stand", "apply", "analyse", "evaluate", or
"create". Here is a short explanation for
the taxonomy, together with verbs that are
commonly used in questions that fit this
taxonomy level:
- remember: recall facts and basic concepts
(verbs: define, duplicate, list, memorize,
repeat, state)
- understand: ...
```

Note also, that we could save on tokens by restricting the prompt to only the one cognitive dimension that a user actually selects in the pipeline. In practice, the instructional part of the prompt is only a small part of the final prompt; the bulk of which consists of the learning materials provided. As such, there is little need to be conservative here, and listing all dimensions might help with clearly delineating them.

Next, we give a brief explanation of the STEx syntax used by the learning objects provided as context:

The learning objects will be given as a list of LaTeX snippets, using the sTeX package for semantic markup. In sTeX, concepts are declared using the or commands, and can be referred to using...

Indeed, we can keep this part brief since the learning objects will naturally contain large amounts of STEx syntax that (ideally) allow the model to pick up on the relevant macros and their usage. In particular, we can safely omit example usages here.

We also make sure that sufficient information is provided in the prompt to ensure the reuse of symbols present in the learning objects and the surrounding modules:

Every learning object will be prefixed with an id that indicates which file it is from, in a format close to what is required in ...

Next, we explain precisely what kinds of questions we expect from the model, namely one of the three autogradable question types supported by STEx (*multiple-choice*, *single-choice* or *fill-in-the-blanks*). We then provide examples for each of them, so the model 'knows' to replicate the relevant STEx macros and environments. We make sure that the examples are representative of the output we expect from the model; that is, they should contain as many detailed annotations as possible, be (in our estimation) interesting and didactically valuable and contain good feedback, especially for the wrong answer options. Therefore, we also point out that adding feedback is desired as follows:

A quiz question can be either a multiple choice question, a single choice question or a fill-in-the-blanks question. We can also provide feedback for each answer option using the 'feedback' key of the relevant macros. This feedback is shown to the student after they have submitted their answer. An example for a multiple choice question is the following; note the semantic markup:

```
``
\begin{sproblem}
  \usemodule[smglom/sets]{mod?bijective}
  \usemodule[smglom/sets]{mod?relation
-composition}
  \usemodule[smglom/arithmetics]{mod?na-
tarith}
  \objective{understand}{bijective}
  \objective{understand}{injective}
  \objective{understand}{surjective}
  Assume  $\text{\fun{f,g}}\backslash\text{NaturalNumbers}$ 
NaturalNumbers$. Which of the following are
true?
  \begin{mcb}
    \mcc[F,feedback={No,  $f$  and  $g$  are
unrelated}]
    {If  $f$  is  $\text{\sn{injective}}$ , so is
```

```
 $g$ }.}
    \mcc[F,feedback={No. since  $f$  does
not need to be  $\text{\sn{surjective}}$ , the
 $\text{\sr{surjective}}\{\text{surjectivity}\}$  of
 $g$  is not enough to make the
 $\text{\sr{compose}}\{\text{composition}\}$  of  $f$ 
and  $g$   $\text{\sn{surjective}}$ .}]
    {If  $f$  is  $\text{\sn{injective}}$  and  $g$ 
is  $\text{\sn{surjective}}$ , then
 $\text{\compose{g,f}}$  is  $\text{\sn{surjec-
tive}}$ .}]
    \mcc[F,feedback={No. Since  $f$  need
not be  $\text{\sn{surjective}}$ , the
 $\text{\sn{composition}}$  need not be  $\text{\sn{-
surjective}}$  either.}]
    {If  $f$  is  $\text{\sn{injective}}$  and  $g$ 
is  $\text{\sn{surjective}}$ , then
 $\text{\compose{g,f}}$  is  $\text{\sn{bijec-
tive}}$ .}]
    \mcc[T]{If  $f$  and  $g$  are  $\text{\sn{injec-
tive}}$ , so is  $\text{\compose{g,f}}$ .}
    \mcc[T]{If  $f$  and  $g$  are  $\text{\sn{sur-
jective}}$ , so is  $\text{\compose{g,f}}$ .}
  \end{mcb}
\end{sproblem}
``
The question is answered correctly, if the
student selects exactly the \mcc
options marked with [T].
```

We delimit the STEx code example using the three backticks "", as is standard in, for example, Markdown. In particular, it is the same delimitation used by GPT (in particular ChatGPT) when producing code in its output.

Furthermore, we deliberately chose an example of a *multiple-choice* question with more than one correct answer—like (in our experience) human beings, GPT too seems to prefer having a single correct choice; in both cases likely because it is easier to come up with questions where there is one correct answer. To direct the model to generate more question types than just 'single choice questions disguised as multiple choice', we promote a variety of question formats.

We subsequently provide similar examples for *single-choice* and *fill-in-the-blank* questions. Finally, we reiterate the expected outcomes and add specific criteria we want the questions to satisfy, starting with ones that seemed important to us from the start:

Once you are given the data described above, you are to reply with the given number of quiz questions as LaTeX code with semantic markup.

Make sure that the questions you generate satisfy the following criteria:

- The answer to a question should give a good indication regarding the extent to which the student has mastered the concept with respect to the given cognitive dimension.
- The questions should be as diverse as possible, i.e. they should test for different

aspects of the concept.

- Make sure you provide good feedback, especially for wrong answers, so that students can learn what they did wrong.
- Assume the students know nothing about the concept other than what's in the learning objects provided.
- In single/multiple choice questions, make sure that the distractors are superficially plausible to students who have not yet mastered the subject, as to not make the problems too easy.
- Feel free to use examples the students might know from elsewhere, e.g. basic high school level knowledge, or foundational principles of the study program the course is in.

Finally, we list the parameters provided by the user as replacement variables that are substituted by the system:

```
%%%%
concepts: %%CONCEPT%%
course: %%COURSE%%
...
learning objects: %%LEARNING_OBJECTS%%
```

4.4.1 | Notable Updates During Iteration

While some aspects of the prompt discussed so far have changed during iterating, the 'basic outline' of the prompt seemed to work surprisingly well and has thus remained largely stable. In response to mistakes we noticed, we added additional criteria to the list at the end, which we discuss now.

Occasionally, there was an excessive tendency in the output to refer to the course materials directly, resulting in generated questions about, for example, specific names of variables in some examples somewhere in the course notes. Hence, we attempt to avoid questions in the output about irrelevant specifics of the course materials:

- We do not want students to rote-memorize definitions, examples, or other text in the learning objects, so never make the correct answer dependent on such details (e.g. variable names, particular examples, etc.).

It also sometimes added feedback or explanations *to the text of the answer choice*, inevitably revealing the correct and wrong answers in the questions itself:

- Do not put any text in the LaTeX code that directly states which answer is correct
- the sTeX macros used above take optional arguments explicitly for that purpose.

In some instances, quiz questions were generated to ask for additional *free-text* answers, so we emphasise that students will not be able to do so:

- Note that students are limited to replying to a question in the form the question type is posed in, i.e. ticking boxes in single/multiple choice questions, or filling in a short text in fill-in-the-blanks questions. They can not provide any additional text.

The output rarely contained fill-in-the-blank questions. When we modify the prompt to explicitly encourage those, problems quickly emerge, which we attempted to fix thusly:

- The correct answer must be unambiguous, particularly for fill-in-the-blanks questions.
- Importantly, the evaluation of student's answers in \fillinsol is done automatically via string matching, so \fillinsol can only contain plain text, no LaTeX code, and students need to type in the answer exactly to get any points.

This seemed to help somewhat, but the primary noticeable result was even fewer fill-in-the-blank questions.

This should not be entirely surprising: beyond the `remember` dimension, fill-in-the-blank questions (where answers are literally string matched to a reference solution) are primarily helpful to ask for an answer that (ideally) can only result from, for example, correctly applying an algorithm, computing some expression and counting specific properties—that is, they are very useful to make questions asking for *complex reasoning* and *processes* autogradable. These questions tend to be particularly difficult to design, though. Hence, the model strongly prefers multiple- and single-choice questions.

In general, the generated questions for the `apply` dimension were not satisfactory. One strategy we tried to improve this was to explicitly point out the use of fill-in-the-blank questions in the prompt, but this did not lead to any noticeable improvement.

4.5 | Evaluation

In the literature, there are two established techniques to evaluate the quality of generated questions (Mulla and Gharpure 2023): (1) automatic evaluation and (2) human-based evaluation.

Certain criteria outlined in Section 3 necessitate specific insights from the course context, accessible only to individuals within the lecture environment. Similarly, assessing the quality of STEx annotations requires expertise and cannot be effectively automated. Hence, an expert-based evaluation was selected—the most common evaluation approach in the context of AQG (Mulla and Gharpure 2023).

For this purpose, we (1) developed a structured survey to evaluate the quality of generated questions based on the criteria above, (2) generated a collection of 30 questions, (3) conducted the survey with experts within the course environment and (4) conducted a qualitative analysis of the survey findings.

4.5.1 | Survey Design

The survey's introductory segment supplied experts with context by outlining the parameters used in the question generation prompt, followed by the complete presentation of the generated question (GQ), incorporating any semantic annotations.

A free-text response field was incorporated to evaluate the accuracy of the content and detect any potential errors in the GQ. This section allowed experts to pinpoint specific content discrepancies or inaccuracies within the GQ. Additionally, we employed a 5-point Likert scale, spanning from 'Very Difficult' to 'Very Easy', to assess the perceived level of difficulty of the GQ.

The survey consisted of six statements to assess the generated question's different quality dimensions. Experts were tasked with indicating their degree of agreement with these statements on a 7-point Likert scale, ranging from 'Strongly Disagree' to 'Strongly Agree'. The statements were outlined as follows:

- The GQ has a good FIT in terms of teaching material.
- The GQ can be solved using the available teaching material.
- The task description of the GQ cannot be misinterpreted (is not ambiguous).
- The GQ is relevant for the achievement of the specified Learning Objective.
- The feedback provided for the answer options of the GQ is helpful.
- The structure of the task corresponds to the specified task format.

Upon completing the survey, the experts could expand their assessments by noting any additional irregularities or remarks. This section, intentionally left open-ended, was aimed at allowing the experts to share extra observations, recommendations or issues they might have about the tasks presented.

4.5.2 | Parameter Selection

To evaluate GQs systematically across multiple topics, we decided on generating a total of 30 questions across six distinct subjects within a university course 'Artificial Intelligence I', focusing on topics in *symbolic AI* with otherwise fixed parameters in the prompt (see Section 4.3).

4.5.2.1 | Difficulty and Cognitive Dimension. The majority of task generation systems documented in the literature focus solely on generating *factual* questions that target the *remember* dimension (Kumar et al. 2023). However, since our main objective in AGQ is to create learning materials capable of effectively enhancing our learner models, it is imperative to include tasks that assess at least the cognitive level of *understand*. Thus, we have opted to confine our investigation in this experiment to tasks that align with the *understand* dimension, assuming that the remaining dimensions are even more challenging to design questions for in our context (as for humans, so

for LLMs). We fixed the *difficulty level* to be 'medium' since we had already determined the model to largely generate relatively easy questions, which are consequently less informative when assessing students' mastery of a topic.

4.5.2.2 | Number of Questions. There is a trade-off regarding the number of questions generated *per prompt* at once: The more questions generated, the more *diverse* they can be expected to be. In contrast, generating fewer per prompt and prompting multiple times prevents immediate consideration of previously generated questions, making it difficult to minimise overlap between questions. However, the model often refuses to generate too many questions at once; it only replies with a few and instructs the user to prompt to yield more. The maximal number of GQs per prompt seems to depend on the topic or, presumably, the amount of course materials added to the prompt, but up to five questions at once seems to work consistently. We leave a more systematic investigation of this trade-off and the resulting quality of *sets* of questions *as a whole* to future work and fixed the number to be five GQs per topic (for a total of 30 GQs).

4.5.2.3 | Topics. We chose six topics from the course spread roughly uniformly across the table of contents of the course notes, with an attempt of making some 'broader' and others more 'specific'; they are given as follows:

1. Arc Consistency (constraint satisfaction problems),
2. Alpha-Beta Search (game play),
3. Semantics of Propositional Logic,
4. Syntax of First-Order Logic,
5. The STRIPS model (planning),
6. The *delete* Relaxation (a heuristic in solving planning problems).

5 | Results and Discussion

5.1 | Question Quality

The **question types** of the GQs were dominated by *single-choice* questions (12 out of 30), followed by *multiple-choice* questions (18 out of 30). The first consistently had three possible answers, while the latter almost always offered four options, usually with three distractors. The constant number of answers may be because the examples in the prompt also contain precisely this number of answers. With the MCQ, explicit care was taken to ensure that the example provided contained more than one correct answer. Nevertheless, the generated MCQs almost exclusively contained exactly three distractors. Notably, no fill-in-the-blank (FIB) questions appear in the generated questions. This is likely because designing meaningful FIB questions is significantly more challenging for the *understand* dimension, especially considering the necessary constraint that the answer be unambiguous.

The **FIT** of the GQs to the provided teaching material was evaluated by the experts as consistently given (28 out of 30),

and most of them were assessed as **solvable** using the available teaching material (27 out of 30). In most cases, the GQs were evaluated as **clearly** and **unambiguously** formulated. Finally, most experts rated the GQs as relevant for achieving the learning objective.

However, the quality of the provided **feedback** was mixed. Frequently, no feedback was generated for an answer option. When it was, it was often not helpful in that it merely rephrased a wrong answer again—that is, the feedback to an incorrect choice X was a variant of ‘No, it is not the case that X ’. This instance of uninformative feedback has also been observed and reported in other studies on LLM-feedback (e.g., (Lohr, Keuning, and Kiesler 2024)).

5.2 | Content Errors

Eleven out of 30 GQs contained errors, mainly in the answer options and the feedback. These errors occurred particularly frequently in the topics *arc consistency* and *semantics of propositional logic*. We can only speculate on the reasons for this. Still, it seems likely that it correlates with the conceptual complexity of the topics or how well the topic is represented in the model’s original training data—semantics requires a much better understanding of the *meaning* of logical statements and the formal/mathematical mechanisms that *provide* statements with meaning (*interpretation functions*, *models* and their formal definitions, etc.), and there are multiple ways to formulate these mathematically. Similarly, to *understand* arc consistency requires a solid intuition about how various constraints in a constraint satisfaction problem *interact* with respect to the variables in a particular problem.

One crucial result to consider is that the model may produce superficially plausible questions that align with the ‘aesthetic’ of a good question but are false in not necessarily trivial ways.

One particularly striking example is the following question on the topic of *propositional logic* (generated in an earlier experiment),

If we have already established $\neg B$, how can we use the Natural Deduction rule for implication ($\Rightarrow$) elimination to infer a new formula?
 $\boxtimes$ Given $B \Rightarrow C$, deduce $\neg C$.

where the answer above was considered the intended and only correct one. This is the common fallacy of *denying the antecedent*, but without understanding the semantics of propositional logic (or the natural deduction calculus) already, this is not immediately obvious.

Notable here is that the question is not only *wrong*, but that it actively *reinforces a common misconception* that may not necessarily be caught even by, for example, student employees who assume they are sufficiently knowledgeable in the domain to evaluate questions, which is arguably even more dangerous than a merely *wrong* question.

5.3 | Semantic Annotations

We can distinguish between two kinds of annotations in STEX: *Structural annotations* and *relational annotations*.

Structural annotations, such as `\begin{sproblem}`, `\begin{mcb}` (for MCQ-blocks), `\mcc[<T|F>`, `feedback={...}]<answer>` (for an answer in a MCQ-block) and `\objective{<competency>}{<symbol>}` should be present in all questions, or all questions of a particular question type, and are very schematic in the sense that they are not particular to a specific question, topic or concept (other than the one explicitly provided as a parameter to the prompt). These annotations seemingly pose no significant challenge to the model—the few examples in the prompt are enough for the model to pick up on them immediately and use them consistently and correctly in (almost) all generated questions. Rarely does it hallucinate plausible but non-existent symbol names in the objectives.

This is *not* the case for **relational annotations**, such as symbol or module references. To annotate a piece of text with a specific symbol, one has to know the module it is declared in, import it (using `\usemodule{...}`) and then annotate the text using the *name* of the symbol (e.g., “`\$a\$ \sr{addition} {plus} \$b\$`” annotates the word *plus* with the symbol named ‘*addition*’ which is currently in scope—globally, there may be multiple symbols with the same name). Naturally, how to do such relational annotations cannot be directly inferred from a few *generic* examples. Instead, humans and AI models need a way to systematically look up symbols and their containing modules if we want to be able to refer to them in annotations.

As mentioned in Section 4.1, we use these annotations to determine the prerequisites of learning objects, including questions; so ideally, we would want the model to annotate as much as possible. Our initial hope that the provided learning objects—which should have ample occurrences of the relevant symbol names and module imports—would be enough for the model to ‘learn’ to use them correctly turned out to be clearly false.⁴ Additional explicit instructions in the prompt also did not improve matters.

We therefore attempted to use the canonical way of giving models the ability to look up additional information: *function calling* and RAG. OpenAI’s API for GPT-4 offers explicit support for providing the model with a set of functions and determining which parameters they need. When prompted, the model will then either return an output as usual, or a JSON object specifying (1) which function the model ‘wants’ to call, and (2) the parameters the call should use. It is then up to the caller of the API to implement functionality that augments the *following* prompt with the correct *return value* for that function call.

We allowed the model to call a *search* function with a sequence of arbitrary strings as parameters. We augmented the prompt by adding a paragraph that instructed the model to use this function to determine how to semantically annotate any text that refers to a concept from the domain of the course. We then used existing functionality to search for any *definitions* in our corpora and

picked the top ten results for each of the string parameters. The search results were subsequently appended to the prompt in a new delimited section, the instruction to call the function was removed (to avoid potentially many subsequent calls or infinite loops), and the resulting prompt was fed back to the model.

This, too, unfortunately, did not have the desired effect; on the contrary. The quality of the questions generated got noticeably worse without improving the quality of the semantic annotations. We conjecture that by appending the search result and making the total prompt significantly longer, individual parts of the prompt have less influence on the output, for example, the instructions at the beginning of the prompt. It is possible that more experimentation, possibly including repeating the instructions multiple times throughout the prompt, can lead to better results in the long run.⁵

5.4 | Limitations

We used a single model (GPT-4-turbo) to generate the questions for the evaluation. At the time of the study (March 2024), this model was—in our opinion—the most suitable model (see Section 4.2 for a detailed explanation). However, we explicitly do not rule out the possibility that there will be ongoing ‘better’ models that are more suitable and can lead to better results. A further limitation of this study is that we did not involve students in evaluating the GQ. To obtain objective evidence of the quality of the tasks, a survey of students on the tasks would be helpful and is planned for the future. Especially when evaluating *multiple-choice* questions, for example, an assessment by (only) experts is not sufficient. Whether a distractor is good or bad usually only becomes apparent when the exercise is presented to students (who try to solve it). If learners never choose the distractor, it is a good indicator that it is too apparent and, therefore, unsuitable. At the time of evaluation, however, there was no possibility of recruiting suitable students from the course. Additionally, most generated questions were not ‘ready to use’ on students without passing a ‘human filter’. This means that we would no longer have evaluated the LLM’s capabilities, but the experts’ ability to filter the generated questions sensibly.

Furthermore, the study did not include any control items. As we did not aim to find out whether the experts were able to distinguish between LLM-generated and human-generated tasks the control items were of less importance. Their only function would have been to normalise the expert’s rating, which was not necessary for our research questions.

6 | Conclusion

In this experiment, we investigated the extent to which LLMs are suitable for generating semantically annotated quiz questions for higher education from semantically annotated course materials that are (1) suitable for a specific course and (2) address *understanding* a concept.

Our research reveals significant limitations in the application of LLMs for this task. While generating questions that address *remembering factual* knowledge works well, creating questions that address *understanding* remains a notable challenge—which would need questions aiming for *conceptual* knowledge or asking

for explanations rather than recall of a simple concept. Despite LLMs’ ability to generate a range of questions, domain experts’ need for extensive filtering underscores the models’ inadequacies in autonomously generating educationally valuable content. The quality of questions generated by LLMs often does not meet educational standards, primarily due to the simplistic nature of the questions and the lack of valuable (and correct) feedback. Thus, the required careful review process by experts is both time-consuming and counter to the goal of automated content creation. In addition, errors in the content of the questions are challenging to detect for non-experts, as semantically incorrect output from LLMs is known to be very close to the ‘truth’ (Sobieszek and Price 2022).

Our results show that generating *structural*, semantic annotations works well. However, this success was not reflected in the case of *relational* annotations, which exhibited poor integration, indicating a limitation in the LLM’s ‘ability’ to contextualise and link concepts within the generated content effectively.

The results of the present work make clear that while LLMs can contribute to the pool of learning materials, their current state requires significant human intervention to refine and validate the generated content. So, despite the promises systems using LLMs give, the human-in-the-loop remains crucial. Although the prompt was carefully designed and we provided a lot of static and dynamic context, the results still show weaknesses in quality and semantic annotation. Nevertheless, further studies with a high number of generated tasks and more experts might show that the automatic generation of questions following our approach could lead to more efficiency when used as an assistive system for the preparation process.

Author Contributions

Dominic Lohr: conceptualization, methodology, data curation, validation, writing – original draft, formal analysis, investigation. **Dennis Müller:** conceptualization, methodology, software, data curation, writing – review and editing, writing – original draft, investigation, supervision. **Marc Berges:** writing – review and editing, methodology, funding acquisition, supervision. **Abhishek Chugh:** software, investigation. **Michael Kohlhasse:** writing – review and editing, funding acquisition, supervision.

Acknowledgements

The work reported in this article was conducted as part of the VoLL-KI (see <https://voll-ki.de>) funded by the Federal Ministry of Education and Research Germany under grant 16DHBKI089. We also would like to thank Rakesh Kumar for the work in developing a significant portion of the front end for the question generator.

Conflicts of Interest

The authors declare no conflicts of interest.

Data Availability Statement

The data that support the findings of this study are available on request from the corresponding author. The data are not publicly available due to privacy or ethical restrictions.

Endnotes

¹ See, for example, <https://toloka.ai/blog/llm-leaderboard/> for a regularly updated comparison across several dimensions.

²It should be mentioned, that the precise meaning of “context window” is not necessarily clear. There is some speculation that the huge context window advertised for GPT-4-Turbo involves neural compression and similar (potentially lossy) techniques to allow for more input rather than representing the actual size of the underlying transformer model’s input layer.

³A common problem is people actively crawling the internet for API keys and publicly accessible APIs to GPT models, intending to abuse those for unintended purposes (e.g., via “prompt hacking”).

⁴One reason for that could be the unfortunately inconsistent usage of annotations in the course materials in general, improvements of which happen frequently but are far from exhaustive.

⁵According to anecdotal “LLM-folklore”, repeating important instructions in a long prompt is supposed to improve results.

References

- Anderson, L. W., and D. R. Krathwohl, eds. 2001. *A Taxonomy for Learning, Teaching, and Assessing: A Revision of Bloom’s Taxonomy of Educational Objectives*. Complete ed. New York: Longman.
- Balse, R., B. Valaboju, S. Singhal, J. M. Warriem, and P. Prasad. 2023. “Investigating the Potential of GPT-3 in Providing Feedback for Programming Assessments.” In *Proceedings of the 2023 Conference on Innovation and Technology in Computer Science Education V.1*, 292–298. Turku, Finland: ACM.
- Bhat, S., H. Nguyen, S. Moore, J. Stamper, M. Sakr, and E. Nyberg. 2022. “Towards Automated Generation and Evaluation of Questions in Educational Domains.” In *Proceedings of the 15th International Conference on Educational Data Mining International Educational Data Mining Society*, edited by A. Mitrovic and N. Bosch, 701–704. Durham, UK: University of Durham.
- Bloom, B. S. 1968. “Learning of Mastery.” *Evaluation Comment* 1, no. 2: 1–12.
- Ch, D. R., and S. K. Saha. 2020. “Automatic Multiple Choice Question Generation From Text: A Survey.” *IEEE Transactions on Learning Technologies* 13, no. 1: 14–25.
- Denny, P., H. Khosravi, A. Hellas, J. Leinonen, and S. Sarsa. 2023. “Can We Trust AI-Generated Educational Content? Comparative Analysis of Human and AI-Generated Learning Resources.” Eprint.
- Dijkstra, R., Z. Genç, S. Kayal, and J. Kamps. 2022. “Reading Comprehension Quiz Generation Using Generative Pre-Trained Transformers.” In *Proceedings of the Fourth International Workshop on Intelligent Textbooks 2022*, edited by S. A. Sosnovsky, P. Brusilovsky, and A. S. Lan, vol. 4, 4–17. Durham, UK: Aachen: CEUR-WS.
- Hattie, J., and H. Timperley. 2007. “The Power of Feedback.” *Review of Educational Research* 77, no. 1: 81–112.
- Hattie, J. A. C. 2009. *Visible Learning: A Synthesis of Over 800 Meta-Analyses Relating to Achievement*. New York: Routledge.
- Hellas, A., J. Leinonen, S. Sarsa, C. Koutchme, L. Kujanpää, and J. Sorva. 2023. “Exploring the Responses of Large Language Models to Beginner Programmers’ Help Requests.” In *Proceedings of the 2023 ACM Conference on International Computing Education Research V.1*, edited by K. Fisler, P. Denny, D. Franklin, and M. Hamilton, 93–105. Chicago, IL, USA: ACM.
- Keller, F. S. 1968. “Good-Bye, Teacher....” *Journal of Applied Behavior Analysis* 1, no. 1: 79–89.
- Kiesler, N., D. Lohr, and H. Keuning. 2023. “Exploring the Potential of Large Language Models to Generate Formative Programming Feedback.” In *2023 IEEE Frontiers in Education Conference (FIE)*, 1–5. College Station, TX, USA: IEEE.
- Kirby, J. R., and M. J. Lawson, eds. 2012. *Enhancing the Quality of Learning: Dispositions, Instruction, and Learning Processes*. Cambridge: Cambridge University Press.
- Kohlhase, M. 2006. “OMDoc—An Open Markup Format for Mathematical Documents [Version 1.2]. No. 4180 in LNAI, Springer Verlag.” <http://omdoc.org/pubs/omdoc1.2.pdf>.
- Kohlhase, M., and D. Müller. 2022. “System Description: sTeX3—A LATEX-Based Ecosystem for Semantic/Active Mathematical Documents.” In *Intelligent Computer Mathematics (CICM) 2022*, edited by K. Buzzard and T. Kutsia, vol. 13467, 184–188. Cham: LNAI Springer. <https://kwarc.info/people/dmueller/pubs/cicm22stexsd.pdf>.
- Kohlhase, M., and D. Müller. n.d. “The sTeX3 Manual.” <https://github.com/slatex/sTeX/blob/main/doc/stex-manual.pdf>.
- Kruse, T., M. Berges, J. Betzendahl, M. Kohlhase, D. Lohr, and D. Müller. 2023. “Learning With ALeA: Tailored Experiences Through Annotated Course Material.” [Object Object]. <https://dl.gi.de/items/af606b9f-a6c9-4106-be59-3b4a89d53f79>.
- Kumar, A. P., A. Nayak, K. Manjula Shenoy, Chaitanya, and K. Ghosh. 2023. “A Novel Framework for the Generation of Multiple Choice Question Stems Using Semantic and Machine-Learning Techniques.” *International Journal of Artificial Intelligence in Education* 34: 332–375.
- Kurdi, G., J. Leo, B. Parsia, U. Sattler, and S. Al-Emari. 2020. “A Systematic Review of Automatic Question Generation for Educational Purposes.” *International Journal of Artificial Intelligence in Education* 30, no. 1: 121–204.
- Lohr, D., M. Berges, M. Kohlhase, D. Müller, and M. Rapp. 2023. “The Y-Model - Formalization of Computer Science Tasks in the Context of Adaptive Learning Systems.” *2023 IEEE 2nd German Education Conference (GECon)*, 1–6. <https://doi.org/10.1109/gecon58119.2023.10295148>.
- Lohr, D., H. Keuning, and N. Kiesler. 2024. “You’re (Not) My Type—Can LLMs Generate Feedback of Specific Types for Introductory Programming Tasks?” *Journal of Computer Assisted Learning*. <https://doi.org/10.1111/jcal.13107>.
- Lohr, D., N. Kiesler, H. Keuning, and J. Jeuring. 2024. “Let Them Try to Figure It Out First - Reasons Why Experts (Do Not) Provide Feedback to Novice Programmers.” *Proceedings of the 2024 on Innovation and Technology in Computer Science Education V 1*: 38–44. <https://doi.org/10.1145/3649217.3653530>.
- Matelsky, J. K., F. Parodi, T. Liu, R. D. Lange, and K. P. Kording. 2023. “A Large Language Model-Assisted Education Tool to Provide Feedback on Open-Ended Responses.” Eprint. <https://arxiv.org/abs/2308.02439>.
- McNichols, H., W. Feng, J. Lee, et al. 2023. “Exploring Automated Distractor and Feedback Generation for Math Multiple-Choice Questions via in-Context Learning.” Eprint. <https://arxiv.org/abs/2308.03234>.
- Mulla, N., and P. Gharpure. 2023. “Automatic Question Generation: A Review of Methodologies, Datasets, Evaluation Metrics, and Applications.” *Progress in Artificial Intelligence* 12, no. 1: 1–32.
- Müller, D. 2023. “An HTML/CSS Schema for TeX Primitives—Generating High-Quality Responsive HTML From Generic TeX.” *TUGboat; TUG 2023 Conference Proceedings*, Volume 44, 275–286. <https://kwarc.info/people/dmueller/pubs/tug23.pdf>.
- OpenAI. 2023. “GPT-4 Technical Report.” *arXiv*. <https://doi.org/10.48550/arXiv.2303.08774>.
- OpenAI. 2024. “OpenAI API Documentation: Prompt Engineering.” <https://platform.openai.com/docs/guides/prompt-engineering>.
- Prather, J., P. Denny, J. Leinonen, et al. 2023. “The Robots Are Here: Navigating the Generative AI Revolution in Computing Education.” In *Proceedings of the 2023 Working Group Reports on Innovation and Technology in Computer Science Education Turku Finland*, 108–159. New York, NY: ACM.
- Sarsa, S., P. Denny, A. Hellas, and J. Leinonen. 2022. “Automatic Generation of Programming Exercises and Code Explanations Using Large Language Models.” In *Proceedings of the 2022 ACM Conference*

on *International Computing Education Research—Volume 1 Lugano and Virtual Event Switzerland*, edited by J. Vahrenhold and K. Fisler, 27–43. New York, NY: ACM.

Sobieszek, A., and T. Price. 2022. “Playing Games With Ais: The Limits of GPT-3 and Similar Large Language Models.” *Minds and Machines* 32, no. 2: 341–364.

Song, T., Q. Tian, Y. Xiao, and S. Liu. 2024. “Automatic Generation of Multiple-Choice Questions for CS0 and CS1 Curricula Using Large Language Models.” In *Computer Science and Education. Computer Science and Technology*, edited by W. Hong and G. Kanaparan, vol. 2023, 314–324. Singapore: Springer Nature Singapore.

Sweller, J. 1988. “Cognitive Load During Problem Solving: Effects on Learning.” *Cognitive Science* 12, no. 2: 257–285.

Sweller, J., J. G. van Merriënboer, and F. W. C. Paas. 1998. “Cognitive Architecture and Instructional Design.” *Educational Psychology Review* 10, no. 3: 251–296.

Touvron, H., T. Lavril, G. Izacard, et al. 2023. “LLaMA: Open and Efficient Foundation Language Models.” <https://arxiv.org/abs/2302.13971>.

Tran, A., K. Angelikas, E. Rama, C. Okechukwu, D. H. Smith, and S. MacNeil. 2023. “Generating Multiple Choice Questions for Computing Courses Using Large Language Models.” In *2023 IEEE Frontiers in Education Conference (FIE)*, 1–8. College Station, TX, USA: IEEE.

Wang, Z., J. Valdez, D. Basu Mallick, and R. G. Baraniuk. 2022. “Towards Human-Like Educational Question Generation With Large Language Models.” In *Artificial Intelligence in Education*, edited by M. M. Rodrigo, N. Matsuda, A. I. Cristea, and V. Dimitrova, vol. 13355, 153–166. Cham: Springer International Publishing.

Wolfe, J. H. 1976. “Automatic Question Generation From Text—An Aid to Independent Study.” In *Proceedings of the ACM SIGCSE-SIGCUE Technical Symposium on Computer Science and Education—Not Known*, edited by R. Colman and P. Lorton, 104–112. New York, NY: ACM Press.

Wong, L. H., H. Park, and C. K. Looi. 2024. “From Hype to Insight: Exploring ChatGPT’s Early Footprint in Education via Altmetrics and Bibliometrics.” *Journal of Computer Assisted Learning* 40: 1428–1446.

Yan, L., L. Sha, L. Zhao, et al. 2024. “Practical and Ethical Challenges of Large Language Models in Education: A Systematic Scoping Review.” *British Journal of Educational Technology* 55, no. 1: 90–112.

Zamfirescu-Pereira, J. D., R. Y. Wong, B. Hartmann, and Q. Yang. 2023. “Why Johnny Can’t Prompt: How Non-AI Experts Try (And Fail) to Design LLM Prompts.” In *Proceedings of the 2023 CHI Conference on Human Factors in Computing Systems CHI’ 23*, edited by A. Schmidt and K. Väänänen. New York, NY, USA: Association for Computing Machinery. <https://doi.org/10.1145/3544548.3581388>.