

Reframing teacher digital competence for generative AI: a task–capability analysis of multi-platform educational discourse

Na Zhao

School of Economics, Harbin University of Commerce, Harbin, China

zhaona@hit.edu.cn

Abstract. Generative Artificial Intelligence (GenAI) has brought the issue of teacher digital competence into focus again, although it is possible to see that most of the frameworks list the skills and do not make any distinction between the tasks of teaching where AI applications are used. The research was based on six platforms provided to determine the expressions of five capabilities in which teachers were able to perform tasks. There were 12,861 records included in the analytic sample and they came from 300 source URLs. The highest capability expressions were efficiency (41.87%) and content development (34.76%). Positive associations of all five task areas were found in pedagogical assistance; the greatest estimates were made with classroom enactment and professional learning. Task-replacement language was less represented in classroom enactment and professional learning after multiplicity adjustment. The results confirm the competence model, which is task sensitive and is based on assessment of capability, orchestration of instruction, the verification of evaluation, mediation of the learner and inquiry of the profession. The corpus explains the expectations of discourse instead of the competence of teachers or effects of their instruction.

Keywords: generative artificial intelligence, teacher digital competence, instructional orchestration, professional agency, social media discourse

1. Introduction

General-purpose language models, the institutional licensing of AI, its use in courses and informal experimentation have all led to the introduction of generative artificial intelligence in education. The competence of teachers is now being offered as a pre-requisite to the educational process of purposeful use instead of being an additional part of the technical provision [1, 2]. The UNESCO approach of recent times has found AI competencies within human-centred agency, ethics, pedagogy, and professional growth. Digital resources are also related with teaching, assessment, learner support, and continuing learning by educators in the same way as DigCompEdu [3, 4]. These paradigms have taken the debate further than the mere utilization of the tools but there remains the need to give an empirical explanation of how the roles of the AI functions are being added on the tasks that the teacher does.

It has been recognized that teacher digital competence is a professional capability. It is based on the situation. The teacher digital competency of Falloon integrates both technical, pedagogical, personal and professional aspects [5]. There are also findings in the literature on teacher education to indicate that competence can be developed as part of programme integration, organisational learning, and technological use within the discipline [6, 7]. TPACK is an enduring description of the interrelationship between technology, teaching and content [8], and Intelligent-TPACK incorporates the teachers moral judgment on the AI decisions made by them [9]. In addition to the researches into AI applied to teachers, there have been systematic studies of practice in the planning, implementation, evaluation, and involvement in the system development [10]. These reports imply that the effective model of competence must identify the role played by AI and the situations where the teacher would make judgments.

The other required strand is AI literacy research. Long and Magerko explain competencies in recognising, utilisation and critical evaluation of AI systems [11], whereas Ng et al. structure the concept of AI literacy around understanding, application, evaluation, creation, and ethics [12]. Hybrid human-AI scholarship puts these abilities into a system of joint monitoring and control [13]. The study of professional agency introduces a time and contextual aspect: teachers rely on previous experience, orient their work toward educational aims, and operate under material and institutional circumstances [14, 15]. The key issue, hence, is alignment. An instructor has to evaluate a competency, locate it in an instructional process, confirm the outcome evidence, facilitate its application with learners and adjust practice by means of professional inquiry.

The majority of empirical researches are based on the measurement of attitudes, self-perceived preparedness or local implementation. The benefits and the drawbacks of the GenAI literature also cover policy issues and uneven quality in methodology [16-19]. The studies on artificial intelligence in learning have been numerous to suggest that the attention should be paid to the teachers, their role in teaching, as well as human accountability [20, 21]. There is another type of evidence that is the public educational discourse. It does not evaluate competence but it will register the salience of AI tasks with specific teaching activities. That semantic organisation can be analyzed with care and can indicate a discursive demand pattern which can be later tested by teacher-development programs using performance tests and classroom experiments.

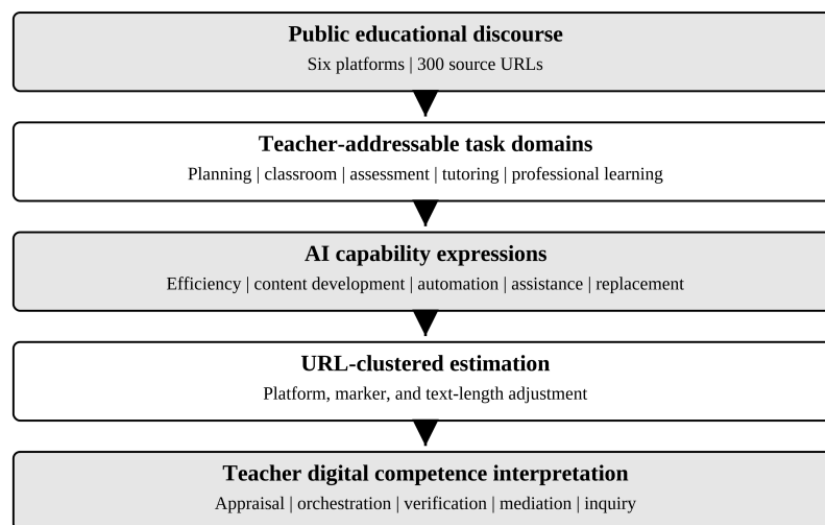


Figure 1. Analytical pathway from public educational discourse to a task-sensitive interpretation of teacher digital competence

The current research focuses on a multi-platform corpus at the source-URL clusters level. Figure 1 is a summary of the analytical path. The paper poses three questions: (RQ1) How are five expressions of AI capabilities spread across teacher-addressable task domains? (RQ2) What capability expressions have task-specific associations that persist even when platform, explicit teacher and AI markers, text length, and URL clustering are adjusted? (RQ3) Which competence architecture can be developed based on these associations without considering discourse as an indicator of teachers actual performance?

2. Materials and methods

2.1. Study design and corpus

The project team provided a combined workbook described as compiled from publicly available pages at Bilibili, Douyin, Kuaishou, Weibo, Xiaohongshu and Zhihu. No new web collection was undertaken by the present study, and no contact was made with platform users. The workbook contained 35,500 substantive records after 12 duplicate header rows were discarded. A main corpus of 32,500 records was distinguished from 3,000 records which had been enriched for risk language in an earlier project workflow through deliberate means. The latter records were excluded since the present analysis concerns teachers tasks and AI capability expressions, not the composition of risk discourse.

The analytic sample had a record that was added when there was at least one teacher-related task term in the recorded educational context field, as well as an AI capability term in the recorded functional field. It resulted in 12,861 entries into all the 300 URLs. There was no full machine-readable acquisition log within the project workbook that would allow the independent recreation of all the steps of collection. Platform labels, URLs, and keyword fields are thus considered as corpus metadata; it does not give an estimate of the prevalence of platforms by users, change over time or opinion of the general population [22, 23].

2.2. Task and capability construction

There were five expressions of capabilities: resource efficiency, content development, automation of processes, pedagogical help and replacement of tasks. The English glosses that have been used to document each category are reported in the supplementary material. All the indicators were binary and non-exclusive, since a record may be having more than one task or capability. The explicit teacher marker in the analytic file was also noted in the case when the text included the standard Chinese words of the teacher, and the explicit AI marker in the A-category field of the workbook. The supplied text was measured as the number of characters. These variables were incorporated into the model to isolate the task capability correlation with the mere difference between the naming of an actor by name, naming of the AI by name, and length of the record.

2.3. Statistical analysis

The paper has been estimated with five binary logistic Generalized Estimating Equation (GEE) models, which were the same as each of the expressions of capability. The clustering unit was the source URL, since it is possible to have a number of records on the same page. All the models had an exchangeable working correlation and robust standard errors [24]. Bilibili was the platform reference category. Five task indicators, five platforms, the explicit teacher marker, the explicit AI marker and standardized text length were in the model:

$$\text{logit}\{Pr(Y_{ij}^{(k)} = 1)\} = \beta_0^{(k)} + \sum_{m=1}^5 \beta_m^{(k)} T_{m,ij} + \sum_{p=1}^5 \gamma_p^{(k)} P_{p,ij} + \delta_1^{(k)} E_{ij} + \delta_2^{(k)} A_{ij} + \delta_3^{(k)} L_{ij}. \quad (1)$$

Y in Equation (1) is the expression of capability k , which is record i of source URL j ; T represents the five task measures; P indicates the non-reference platform measures; E and A are the explicit teacher and AI signs; and L stands for the standard length of the text. The findings were presented with adjusted Odds Ratio (ORs), 95% Confidence Intervals (CIs), as well as joint Wald tests. There was a Benjamini-Hochberg correction on the p values of the 25 tasks-coefficient. Additional analysis included an independent working correlation, uniform weighting of URLs, mutually exclusive one-task subset, leave-one-platform-out refits, pedagogical assistance and replacement of the task.

2.4. Data protection

The analytic workbook had no usernames, account identifiers, private messages, geolocation or profile attributes. The paper does not reproduce any verbatim post and it does not distribute raw text. These decisions are based on data minimization, purpose limitation, de-identification, non-reidentification, and contextual integrity principles that are formulated in internet-research guidance [25-27]. The handling protocol was developed with regard to the Personal Information Protection Law and Data Security Law of the People Republic of China and general principles of GDPR [28-30].

3. Results

3.1. Analytic corpus and semantic distribution

The total number of records in the analytic sample was 12,861. Bilibili contributed 2,878 records, Douyin 2,554, Kuaishou 1,995, Weibo 868, Xiaohongshu 2,200, and Zhihu 2,366. The mean text length was 66.23 characters ($SD = 23.77$). An explicit teacher marker appeared in 6,895 records, and the workbook's explicit AI marker appeared in 6,211 records. Most records contained one task domain ($n = 11,319$); 1,497 contained two domains, and 45 contained three. Capability overlap was more common: 8,861 records contained one capability, 3,419 contained two, 548 contained three, and 33 contained four.

Table 1. Distribution of teacher-addressable task domains and AI capability expressions

Category	Records, n	Record-weighted, %	URL-balanced, %
Task domains			
Instructional planning and resource design	3,048	23.70	24.38
Classroom enactment and interaction	3,368	26.19	25.92
Assessment and feedback	4,263	33.15	32.64
Learner explanation and tutoring	2,334	18.15	18.12
Professional learning and inquiry	1,435	11.16	11.80
AI capability expressions			
Resource efficiency	5,385	41.87	40.51
Content development	4,471	34.76	35.17
Process automation	2,272	17.67	17.37
Pedagogical assistance	3,323	25.84	25.16
Task replacement	2,024	15.74	15.66

It is seen in Table 1 that the most common task area was assessment and feedback (33.15%). The highest capability expressions were resource efficiency (41.87) and content development (34.76). Percentages of

URLs, which were equal to the percentages of records, came after the record-weighted ones, meaning that no group of high volume URLs could have been responsible for the rank order. These numbers are an explanation of the analytic corpus.

3.2. Adjusted task–capability associations

The joint task test was significant for resource efficiency, $\chi^2(5) = 13.32$, $p = .021$, although no individual task coefficient remained reliable after multiplicity adjustment. The joint test for content development approached conventional significance, $\chi^2(5) = 10.51$, $p = .062$. Classroom enactment was associated with less content-development language (OR = 0.83, 95% CI [0.73, 0.94], $q = .008$). Process automation showed no overall task differentiation, $\chi^2(5) = 5.90$, $p = .316$. These results place efficiency, content development, and automation mainly within a general productivity layer rather than a sharply segmented task structure.

Pedagogical assistance showed the clearest task pattern, $\chi^2(5) = 91.81$, $p < .001$. Every task domain retained a positive association after multiplicity adjustment. Classroom enactment produced the largest estimate (OR = 1.93, 95% CI [1.67, 2.24]), followed by professional learning (OR = 1.82, 95% CI [1.47, 2.26]), planning (OR = 1.76, 95% CI [1.54, 2.02]), tutoring (OR = 1.62, 95% CI [1.36, 1.94]), and assessment (OR = 1.61, 95% CI [1.39, 1.85]). The discourse thus connected assistance with the full span of teacher-addressable work rather than with a single routine.

Task-replacement language was also different in the task, $\chi^2(5) = 13.53$, $p = .019$. Classroom enactment was negatively associated with replacement (OR = 0.79, 95% CI [0.66, 0.94], $q = .026$), and professional learning showed the strongest negative association (OR = 0.70, 95% CI [0.57, 0.86], $q = .002$). The assessment coefficient was negative at the unadjusted level (OR = 0.82, 95% CI [0.68, 0.99], $p = .037$) but did not survive the 25-test adjustment ($q = .103$). Table 2 gives all the coefficients of the tasks.

Table 2. Adjusted task associations with AI capability expressions

Task domain	Resource efficiency	Content development	Process automation	Pedagogical assistance	Task replacement
Planning and resource design	0.99 [0.88, 1.13]	0.89 [0.77, 1.02]	1.00 [0.81, 1.23]	1.76 [1.54, 2.02]	0.84 [0.71, 1.00]
Classroom enactment	0.96 [0.85, 1.08]	0.83 [0.73, 0.93]	1.00 [0.84, 1.20]	1.93 [1.67, 2.24]	0.79 [0.66, 0.94]
Assessment and feedback	1.08 [0.95, 1.23]	0.88 [0.77, 1.01]	1.02 [0.84, 1.23]	1.61 [1.39, 1.85]	0.82 [0.68, 0.99]
Explanation and tutoring	1.08 [0.94, 1.23]	0.94 [0.82, 1.08]	0.86 [0.71, 1.04]	1.62 [1.36, 1.94]	0.85 [0.70, 1.02]
Professional learning	0.90 [0.76, 1.07]	0.92 [0.79, 1.07]	1.01 [0.85, 1.22]	1.82 [1.47, 2.26]	0.70 [0.57, 0.86]

Note. The ORs are adjusted by the cells and the 95% CI is reported. Platform, five task indicators, explicit teacher and AI signs, as well as standardized text length; the standard errors of all models are clustered by source URL. Bold entries have Benjamini–Hochberg $q < 0.05$ across the 25 task coefficients. OR > 1 indicates greater representation of the capability expression within the task domain.

3.3. Sensitivity analysis

All the five exchangeable GEE models were found to have close working correlations, which were estimated at 0.0041 or less in absolute value. Re-estimation with an independent working correlation changed task ORs

by less than 0.007. Equal URL weighting preserved the ordering of the task and capability distributions shown in Table 1. In leave-one-platform-out models, all five pedagogical-assistance estimates remained above 1.0, with OR ranges of 1.55–1.99 for planning, 1.68–2.27 for classroom enactment, 1.43–1.83 for assessment, 1.42–1.82 for tutoring, and 1.25–2.18 for professional learning. The negative association between professional learning and replacement remained below 1.0 in every platform omission (OR range = 0.67–0.75).

4. Discussion

4.1. A generic productivity layer

The corpus focused on resource efficiency and content development, yet the modified models did not find much stable differentiation of these abilities in teacher work. This outcome is significant since competence frameworks may be too detailed with each tool feature being mapped to a different skill. The current data indicate that there exists a general productivity layer: teachers must determine whether a system can minimize routine preparation, change material or automate a process without presupposing that such functions have a fixed pedagogical meaning. The learning worth of an outline created, translated, batch processed or summarized will rely on the objective of the lesson, the disciplinary level and the proof that educators demand of students. Capability appraisal precedes tool operation therefore.

4.2. Assistance as a cross-task pedagogical relation

All teacher-addressable tasks had a positive correlation with pedagogical assistance. The trend was most evident in classroom enactment and professional learning, which are two areas that require interpretation, responsiveness and reflection. An assistance expression may be an aid to a lesson plan, explanation, feedback or inquiry but each application needs a teacher to place the system in an educational activity as it unfolds. The result is in line with hybrid human-AI explanations that maintain professional monitoring and control [13] and teacher-AI studies where teachers have planning, implementation, assessment, and system-evaluation functions [10]. Instructional orchestration (the ability of the teacher to define an educational purpose, select a suitable AI function, establish conditions of use, and update the activity when the output is not satisfactory) should therefore be evaluated in a competence programme.

4.3. Professional judgment and replacement language

The fact that replacement language is less represented in classroom enactment and professional learning means that public educational discourse can differentiate some teacher work with easily replaceable routines. The teaching of the classroom involves situated explanation, social coordination and interpretation of learner responses in real time. Teachers need to look at practice, talk about evidence and change shared norms as part of professional learning. These activities are based on professional agency which is attained by combining personal capacity and cultural, material and institutional conditions [14, 15]. This outcome does not confirm that AI will not be able to do parts of these tasks. It demonstrates that the corpus was more likely to depict AI as an aid rather than a total substitute when the task involved context-sensitive judgment.

The finding also explains the role of evaluative verification. Plausible output cannot be treated as sufficient educational evidence by teachers. They have to verify factual adequacy, disciplinary quality, provenance, attribution and the relationship between a generated product and a learners own work. Current GenAI research repeatedly identifies unreliable output, overreliance, integrity concerns, and policy inconsistency [16-19, 23]. Verification therefore belongs inside instructional design rather than after it. A teacher who can operate a

model but cannot establish evidence conditions has technical fluency without adequate professional competence.

4.4. A task-sensitive competence architecture

Figure 2 and Table 3 combine the empirical pattern with the established competence and agency literature. There are five interrelated fields in the architecture. Capability appraisal is related to the choice and critical questioning of AI functions. Instructional orchestration deals with the integration of these functions with goals, order, involvement and discipline. Evaluative verification relates to the quality and attribution of outputs and learning evidence. Mediation that involves learners is concerned with explanation, feedback, prompting and self-regulation. Professional inquiry involves collective review, adaptation and further development of shared practice. The entire sequence was based on task capability alignment and context sensitive professional agency.

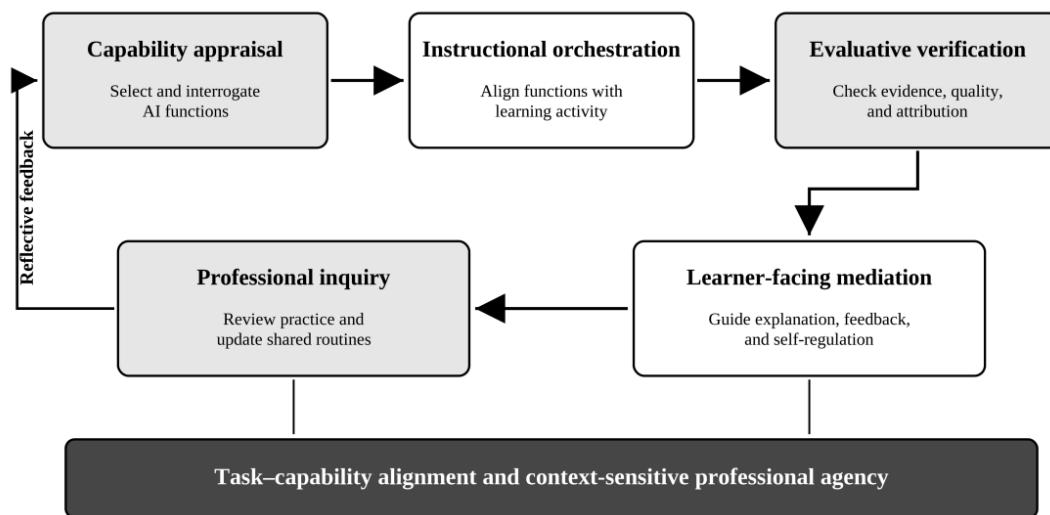


Figure 2. Task-sensitive competence architecture derived from the empirical associations and established teacher-competence scholarship. The figure is a conceptual synthesis, not a validated measurement scale

Table 3. Empirically grounded domains for teacher development and assessment

Competence domain	Empirical anchor	Teacher capability	Observable evidence in development
Capability appraisal	Efficiency and content development dominated across tasks.	Compares functions, limitations, data conditions, and disciplinary fit before use.	Tool-comparison rationale and documented error audit.
Instructional orchestration	Assistance was positively associated with all five tasks.	Positions AI within objectives, sequence, roles, interaction, and contingency plans.	AI-explicit lesson design and microteaching enactment.

Table 3. Continued

Evaluative verification	Replacement was less represented in classroom and professional learning.	Checks accuracy, provenance, attribution, process evidence, and learner authorship.	Annotated verification log and revised assessment rubric.
Learner-facing mediation	Assistance was associated with tutoring and classroom work.	Guides prompting, explanation, feedback interpretation, and learner self-regulation.	Observed feedback conference or tutoring simulation.
Professional inquiry	Professional learning combined high assistance with low replacement.	Uses AI-supported evidence while retaining collective review and professional responsibility.	Peer-reviewed inquiry memo and updated shared protocol.

The design has practical consequences to faculty development. Workshops that are based on product features or prompt recipes should not be organised by institutions. A more robust design would employ real teaching tasks, ask the participants to explain why task-capability alignment is appropriate, analyze the output produced using disciplinary standards and practice learner-facing mediation. Artefact review, observation and reflective explanation must be integrated in assessment instead of self-reported confidence. The same method can assist teachers to convert workplace practices into assessable learning activities in industry-linked and applied curricula. There is no direct measure of enterprise participation in the current corpus though, therefore, such application needs evidence provided independently by teachers, students and industry.

5. Conclusion

The analysis found a stable difference in multi-platform educational discourse. Efficiency, content development and automation were a widely distributed productivity layer while pedagogical support was appended to every teacher addressable task. Replacement language was less visible in classroom enactment and professional learning where the work of education is very much dependent on situated judgment and reflective practice. These results inform a task-sensitive explanation of teacher digital competence that links capability appraisal, instructional orchestration, evaluative verification, learner-facing mediation, and professional inquiry. The suggested architecture does not regard discourse as proof of achieved competence. It provides an arguable framework of future instrument creation, faculty education, and classroom validation.

References

[1] Miao, F., & Cukurova, M. (2024). *AI competency framework for teachers*. UNESCO. <https://doi.org/10.54675/ZJTE2084>

[2] Miao, F., & Holmes, W. (2023). *Guidance for generative AI in education and research*. UNESCO. <https://doi.org/10.54675/EWZM9535>

[3] Redecker, C. (2017). *European framework for the digital competence of educators: DigCompEdu* (Y. Punie, Ed.). Publications Office of the European Union. <https://doi.org/10.2760/159770>

[4] Caena, F., & Redecker, C. (2019). Aligning teacher competence frameworks to 21st century challenges: The case for the European Digital Competence Framework for Educators (DigCompEdu). *European Journal of Education*, 54(3), 356–369. <https://doi.org/10.1111/ejed.12345>

[5] Falloon, G. (2020). From digital literacy to digital competence: The teacher digital competency (TDC) framework. *Educational Technology Research and Development*, 68, 2449–2472. <https://doi.org/10.1007/s11423-020-09767-4>

- [6] Instefjord, E. J., & Munthe, E. (2017). Educating digitally competent teachers: A study of integration of professional digital competence in teacher education. *Teaching and Teacher Education*, 67, 37–45. <https://doi.org/10.1016/j.tate.2017.05.016>
- [7] Pettersson, F. (2018). On the issues of digital competence in educational contexts—A review of literature. *Education and Information Technologies*, 23, 1005–1021. <https://doi.org/10.1007/s10639-017-9649-3>
- [8] Mishra, P., & Koehler, M. J. (2006). Technological pedagogical content knowledge: A framework for teacher knowledge. *Teachers College Record*, 108(6), 1017–1054. <https://doi.org/10.1111/j.1467-9620.2006.00684.x>
- [9] Celik, I. (2023). Towards intelligent-TPACK: An empirical study on teachers' professional knowledge to ethically integrate artificial intelligence (AI)-based tools into education. *Computers in Human Behavior*, 138, Article 107468. <https://doi.org/10.1016/j.chb.2022.107468>
- [10] Celik, I., Dindar, M., Muukkonen, H., & Järvelä, S. (2022). The promises and challenges of artificial intelligence for teachers: A systematic review of research. *TechTrends*, 66, 616–630. <https://doi.org/10.1007/s11528-022-00715-y>
- [11] Long, D., & Magerko, B. (2020). What is AI literacy? Competencies and design considerations. In *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems* (pp. 1–16). Association for Computing Machinery. <https://doi.org/10.1145/3313831.3376727>
- [12] Ng, D. T. K., Leung, J. K. L., Chu, S. K. W., & Qiao, M. S. (2021). Conceptualizing AI literacy: An exploratory review. *Computers and Education: Artificial Intelligence*, 2, Article 100041. <https://doi.org/10.1016/j.caeai.2021.100041>
- [13] Molenaar, I. (2022). Towards hybrid human–AI learning technologies. *European Journal of Education*, 57(4), 632–645. <https://doi.org/10.1111/ejed.12527>
- [14] Eteläpelto, A., Vähäsantanen, K., Hökkä, P., & Paloniemi, S. (2013). What is agency? Conceptualizing professional agency at work. *Educational Research Review*, 10, 45–65. <https://doi.org/10.1016/j.edurev.2013.05.001>
- [15] Biesta, G., Priestley, M., & Robinson, S. (2015). The role of beliefs in teacher agency. *Teachers and Teaching*, 21(6), 624–640. <https://doi.org/10.1080/13540602.2015.1044325>
- [16] Kasneci, E., Sessler, K., Küchemann, S., Bannert, M., Dementieva, D., Fischer, F., Gasser, U., Groh, G., Günnemann, S., Hüllermeier, E., Krusche, S., Kutyniok, G., Michaeli, T., Nerdel, C., Pfeffer, J., Poquet, O., Sailer, M., Schmidt, A., Seidel, T., ... Kasneci, G. (2023). ChatGPT for good? On opportunities and challenges of large language models for education. *Learning and Individual Differences*, 103, Article 102274. <https://doi.org/10.1016/j.lindif.2023.102274>
- [17] Bond, M., Khosravi, H., De Laat, M., Bergdahl, N., Negrea, V., Oxley, E., Pham, P., Chong, S. W., & Siemens, G. (2024). A meta systematic review of artificial intelligence in higher education: A call for increased ethics, collaboration, and rigour. *International Journal of Educational Technology in Higher Education*, 21, Article 4. <https://doi.org/10.1186/s41239-023-00436-z>
- [18] Farrokhnia, M., Banihashem, S. K., Noroozi, O., & Wals, A. (2024). A SWOT analysis of ChatGPT: Implications for educational practice and research. *Innovations in Education and Teaching International*, 61(3), 460–474. <https://doi.org/10.1080/14703297.2023.2195846>
- [19] Chan, C. K. Y. (2023). A comprehensive AI policy education framework for university teaching and learning. *International Journal of Educational Technology in Higher Education*, 20, Article 38. <https://doi.org/10.1186/s41239-023-00408-3>
- [20] Ouyang, F., & Jiao, P. (2021). Artificial intelligence in education: The three paradigms. *Computers and Education: Artificial Intelligence*, 2, Article 100020. <https://doi.org/10.1016/j.caeai.2021.100020>
- [21] Holmes, W., Porayska-Pomsta, K., Holstein, K., Sutherland, E., Baker, T., Buckingham Shum, S., Santos, O. C., Rodrigo, M. T., Cukurova, M., Bittencourt, I. I., & Koedinger, K. R. (2022). Ethics of AI in education: Towards a community-wide framework. *International Journal of Artificial Intelligence in Education*, 32(3), 504–526. <https://doi.org/10.1007/s40593-021-00239-1>

- [22] Hwang, G.-J., Xie, H., Wah, B. W., & Gašević, D. (2020). Vision, challenges, roles and research issues of artificial intelligence in education. *Computers and Education: Artificial Intelligence*, 1, Article 100001. <https://doi.org/10.1016/j.caeai.2020.100001>
- [23] Zawacki-Richter, O., Marín, V. I., Bond, M., & Gouverneur, F. (2019). Systematic review of research on artificial intelligence applications in higher education—Where are the educators? *International Journal of Educational Technology in Higher Education*, 16, Article 39. <https://doi.org/10.1186/s41239-019-0171-0>
- [24] Liang, K.-Y., & Zeger, S. L. (1986). Longitudinal data analysis using generalized linear models. *Biometrika*, 73(1), 13–22. <https://doi.org/10.1093/biomet/73.1.13>
- [25] Franzke, A. S., Bechmann, A., Zimmer, M., Ess, C. M., & Association of Internet Researchers. (2020). *Internet research: Ethical guidelines 3.0*. <https://aoir.org/reports/ethics3.pdf>
- [26] Fiesler, C., & Proferes, N. (2018). “Participant” perceptions of Twitter research ethics. *Social Media + Society*, 4(1), 1–14. <https://doi.org/10.1177/2056305118763366>
- [27] Williams, M. L., Burnap, P., & Sloan, L. (2017). Towards an ethical framework for publishing Twitter data in social research: Taking into account users' views, online context and algorithmic estimation. *Sociology*, 51(6), 1149–1168. <https://doi.org/10.1177/0038038517708140>
- [28] Standing Committee of the National People's Congress. (2021). *Personal information protection law of the People's Republic of China*. https://en.spp.gov.cn/2021-12/29/c_948419.htm
- [29] Standing Committee of the National People's Congress. (2021). *Data security law of the People's Republic of China*. https://www.npc.gov.cn/englishnpc/c2759/c23934/202112/t20211209_385109.html
- [30] European Parliament & Council of the European Union. (2016). Regulation (EU) 2016/679 (General Data Protection Regulation). *Official Journal of the European Union*, L119, 1–88. <https://eur-lex.europa.eu/eli/reg/2016/679/oj/eng>