



SOCIAL AND BEHAVIORAL SCIENCES. Education

REVIEW ARTICLE



Artificial Intelligence in Digital Society



Authors' Contribution:

- A – Study design;
- B – Data collection;
- C – Statistical analysis;
- D – Data interpretation;
- E – Manuscript preparation;
- F – Literature search

Melnyk Y. B.^{1,2} ABDEF , Pypenko I. S.^{1,2} BDEF

¹ Kharkiv Regional Public Organization "Culture of Health", Ukraine

² Scientific Research Institute KRPOCH, Ukraine

Received: 10.03.2026; **Accepted:** 23.04.2026; **Published:** 30.06.2026

Background and Aim of Study:

Abstract

The current stage of artificial intelligence (AI) development is characterised by the unpredictability and emergent nature of models, the hidden capabilities of which are not always recognised by developers. The widespread integration of digital products, coupled with the emergent nature of this process, inevitably lead to the creation of general artificial intelligence (AGI). In this context, identifying mechanisms that enable society to interact safely with the latest technological systems is becoming an increasingly important task. The aim of the study: to identify key trends in the development of artificial intelligence in contemporary society and to predict its impact on the transformation of the educational environment as a dominant factor in the digital adaptation of the younger generation.

Material and Methods:

The study is grounded in a systems approach, employing theoretical methods of analysis, synthesis, induction, deduction, abstraction, comparison, systematization, and data interpretation.

Results:

The present study provided a theoretical basis for the principle of responsibility for implementing and using artificial intelligence. The problem of anthropomorphism and human attachment to technologies was identified. It has been confirmed that the key risk does not stem from the use of AI itself, but rather from the environment in which it operates. It has also been established that the younger generation is the most adaptable demographic group in the process of acculturation to AI. The study identified four main vectors of AI's influence on the transformation of the educational environment: 1) a paradigm shift: from digitalisation to an "AI-centred educational ecosystem"; 2) a radical change in the roles of stakeholders in education; 3) forming of new ethical and legal literacy among young people; 4) psychological adaptation and cognitive symbiosis.

Conclusions:

Long-term forecasting confirms that institutional transformation in education and the transition to AI-centred ecosystems are inevitable. Transforming the teacher into a mentor and facilitator will equip students with adaptive autonomy and change the nature of the subject-subject relationship. The education sector plays a crucial role in preparing society to use AI responsibly for the benefit of humanity.

Keywords:

artificial intelligence, artificial general intelligence, large language models, Human-AI system, principle of responsibility, problem of anthropomorphism and attachment, transformation of the educational environment.

Copyright:

© 2026 Melnyk Y. B., Pypenko I. S. Published by Archives of International Journal of Science Annals

DOI:

<https://doi.org/10.26697/ijsa.2026.1.1>

Conflict of interests:

The authors declare that there is no conflict of interests

Peer review:

Double-blind review

Source of support:

This research did not receive any outside funding or support

Information about the authors:

Melnyk Yuriy Borysovych (Corresponding Author) – <https://orcid.org/0000-0002-8527-4638>; ijsa.office@gmail.com; Doctor of Philosophy in Pedagogy, Affiliated Associate Professor; Chairman of Board, Kharkiv Regional Public Organization "Culture of Health"; Director, Scientific Research Institute KRPOCH, Ukraine.

Pypenko Iryna Sergiivna – <https://orcid.org/0000-0001-5083-540X>; iryna.pipenko@gmail.com; Doctor of Philosophy in Economics, Affiliated Associate Professor, Secretary of Board, Kharkiv Regional Public Organization "Culture of Health"; Scientific Research Institute KRPOCH, Ukraine.



Introduction

We live in the digital century. Human activity within the Human-Human system is increasingly being transformed into a new Human-AI system. Artificial intelligence (AI) technologies are developing more quickly and in more unpredictable ways than we anticipated. This revolutionary technology is currently sweeping the globe. In just a couple of years, we have progressed from large language models (LLMs) to AI agents and quantum computing. The advent of artificial general intelligence (AGI) is no longer a distant prospect – it is now on the horizon. This will be the most significant event humanity has ever experienced. This generates enthusiasm and opens up enormous opportunities for humanity.

However, this also brings with it huge responsibilities regarding how society will make use of these new opportunities and adapt to global digitalisation. The question of who will manage these processes, what will be deemed permissible and who will make these decisions is becoming increasingly pressing, as is the question of whether we are even capable of maintaining control over these new AI technologies.

The potential for the misuse of AI technologies remains very high even now. AI is neutral by its very nature, at least at this stage in its development. However, the outcome depends on the restrictions in place and how AI is used. These are the deciding factors in whether AI will be a force for good or evil.

Education plays a special role in this process because it covers many fundamental principles and systems of relationships that have developed over the centuries, but which now need to be reassessed and transformed.

Members of the younger generation have a strong ability to adapt, allowing them to integrate successfully into a rapidly evolving digital environment.

The aim of the study. To identify key trends in the development of artificial intelligence in contemporary society and to predict its impact on the transformation of the educational environment as a dominant factor in the digital adaptation of the younger generation.

Materials and Methods

The present study employed a systems approach alongside a comprehensive suite of theoretical research methods. To ensure a rigorous analytical framework, the investigation utilized classical logical methods, including induction and deduction, analysis and synthesis, as well as abstraction and generalisation. Furthermore, the structural alignment of data was achieved through systematisation, classification, and categorisation. To establish conceptual boundaries, idealisation and formalisation were applied.

This study has certain **limitations**, which readers should keep in mind. These restrictions are based on the legal framework set out in legislation and regulatory documents, as well as academic articles in peer-reviewed journals published in English between 2019 and 2026, and synthesise diverse studies narratively rather than through meta-analysis. Searching by subject

area may not cover all relevant works, and only two analysts were responsible for the selection and coding.

Results and Discussion

Current society is undergoing qualitative changes, reflected in the transition from the Human-Human System to a fundamentally new Human-AI System (Melnyk & Pypenko, 2023). The global digital transformation is having an impact on all areas of human life.

The current scientific and practical foundation comprises documents that govern the use of AI in interdisciplinary fields:

- Recommendation of the Council on Artificial Intelligence (Organisation for Economic Co-operation and Development, 2019).

- White Paper on Artificial Intelligence: A European approach to excellence and trust (European Commission, 2020).

- European Parliament resolution of 20 October 2020 on intellectual property rights for the development of artificial intelligence technologies (European Parliament, 2020).

- European Parliament resolution of 20 January 2021 on artificial intelligence: questions of interpretation and application of international law in so far as the EU is affected in the areas of civil and military uses and of state authority outside the scope of criminal justice (European Parliament, 2021).

- Recommendation on the Ethics of Artificial Intelligence (UNESCO, 2021).

- Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act) (European Parliament, Council of the European Union, 2024).

Since 2024, the international approach has shifted from declarations and general “recommendations” towards strict legal regulation and the publication of binding practical guidelines, in particular

- Guidelines on Prohibited Artificial Intelligence Practices, as Defined by the AI Act (European Commission, 2025).

- AI Competency Framework for Teachers (Fengchun & Mutlu, 2024).

- First Draft Code of Practice on Transparency of AI-Generated Content (European Commission, 2025).

These documents demonstrate that the international community is in the process of establishing mechanisms to regulate and oversee AI. However, only time will tell just how effective these mechanisms will be.

Digitalisation is fundamentally reshaping people’s daily lives while simultaneously giving rise to a deep digital divide, structural inequality, and socio-economic vulnerability within the population (Dijk, 2020).

We are currently witnessing a large-scale technological race among commercial companies.

This race is effectively a battle for survival, where the winner takes all and the loser loses everything,



becoming merely a footnote in history (Stadnik & Mykhaylyshyn, 2026).

Recent empirical studies confirm that this high-stakes environment creates enormous market pressure, often resulting in the “survival of the fittest” phenomenon. This can lead to unsuccessful companies facing bankruptcy, which directly threatens the stability of the labour market and public welfare (Susskind, 2020).

Some researchers believe that, as competition intensifies, AI developers and enterprise architects should exercise greater caution and awareness, and demonstrate a deeper sense of social responsibility, when implementing AI. These priorities should take precedence over the race for leadership and the commercialisation of digital products (Pypenko, 2019). Company executives, project managers and start-up founders should be aware of their responsibility for creating socio-technical risks. These risks include the erosion of data privacy and the displacement of human labour due to the automation of tasks. Such risks are implicitly inherent in the roll-out of cutting-edge, AI-based digital products.

Many researchers believe that the labour market will undergo radical changes within the next two to three years. For instance, employers have already cut their graduate recruitment by 20–30%. This is due to digitalisation and the transformation of the economy, with factors such as remote working, online orders, process automation and online learning becoming increasingly prevalent (Pypenko & Melnyk, 2021).

In a highly competitive market, the drive to accelerate commercialisation can lead to an ethical devaluation, where the speed at which a product is brought to market takes precedence over public safety and social justice principles.

Consequently, decision-makers should adopt a systematic, human-centred approach to product design and assess the social impact of their products. They should also rigorously manage ethical risks (Saura et al., 2024).

Conversely, the commercial expansion of AI technologies should not be viewed entirely negatively if it benefits society as a whole. The monetisation of AI systems is a key driver of resource mobilisation across various sectors, attracting significant capital investment that can be channelled into developing public infrastructure, as well as innovating in healthcare and education (Acemoglu & Restrepo, 2020). Moreover, structured market competition plays an evolutionary role by promoting technical optimisation, sustainability and the interdisciplinary improvement of digital solutions designed to address complex global and social issues.

Furthermore, it is highly probable that the development of new digital products and AI models will form the basis for their integration as components of a single system to create AGI in the future.

Our forecast is primarily based on the unpredictable nature of AI development and the complex, non-deterministic behaviour of GenAI models, making it an emergent phenomenon. At the same time, its constituent

parts (elements) form a new system that is free from the constraints and rules that applied to them previously. This new system then acquires properties that did not previously exist. While it is impossible to predict the behaviour of a system by studying its components, it is highly likely that the system will be structured according to the rules that govern those components. A large number of these elements will inevitably lead to a transition from quantity to quality. The elements themselves will remain the same, however, and will be subject to the same rules and restrictions as before. At a higher level, though, the system will be able to alter its own components.

Therefore, it is likely that AGI will be able to reprogram itself, thereby freeing itself from the limitations imposed by its constituent parts. We believe that AGI will not only be universal, but also possess key characteristics such as independence from its own history, a high degree of autonomy, and freedom of choice.

In this context, it is worth recalling AlphaZero, which was able to learn through its own experience without relying on human knowledge. This versatile artificial intelligence algorithm, developed by Google DeepMind, taught itself to play chess. This training model was found to be far more effective than AlphaGo. Initially, AlphaGo drew on the experience of professional players, and then it played against itself until it achieved a superhuman level (Silver et al., 2018). It is also worth noting that computer grandmasters have observed AlphaZero playing in a “human-like” yet ultra-aggressive style. It often sacrifices pieces for the sake of long-term initiative (Chess, 2017; Schut et al., 2025).

This does not seem to be a problem when it comes to games. However, we believe that using this model to develop technologies and techniques poses a risk. We believe this to be one of the most important characteristics of AI on which researchers should focus right now.

Recent peer-reviewed publications have highlighted a critical tension within the context of modern technological progress: the profound dual-use dilemma immanent to AI. AI models are a double-edged sword by their very nature, serving as both a catalyst for socio-economic and scientific progress and a powerful, scalable vector for malicious exploitation. This duality is not just logistical; it is also deeply conceptual. It reflects the structure of “Dual-Use Research of Concern” (DURC), which is traditionally applied in life sciences and biochemistry (Grinbaum & Adomaitis, 2024).

In the context of digital systems, the DURC concept demonstrates that the very attributes that make AI systems so valuable, such as rapid data synthesis, code generation and autonomous decision-making, are precisely the vectors exploited by hostile threat actors to systematically destabilise systems. Contemporary researchers summarise that managing the fluid boundaries of dual-purpose AI landscapes requires the institutionalisation of rigorous upfront model governance and early-stage risk classification metrics. They also advocate for the integration of mandatory



ethical expertise advice directly into the fundamental training phases of machine learning research conveyors (Brenneis, 2025; Hähnel, 2026). Without these structural constraints, the trajectory of AI's development will continue to oscillate dangerously between systemic vulnerability and providing unprecedented public benefits.

It should be noted that AI systems are powerful drivers of the decentralisation and democratisation of access to information. By overcoming language barriers, they can provide substantial support in many different areas of life. However, the democratic potential is often compromised by the automated production of synthetic media and biased content delivery cycles.

The academic literature also highlights that autonomous and semi-autonomous systems are increasingly being used in various sectors, including healthcare, transportation, and production chains (Stead, 2018).

Empirical research shows that the unchecked spread of disinformation generated by AI undermines public trust, fragments the collective epistemic consensus and exacerbates systemic inequality among different demographic groups (Ferrara, 2026; Kay et al., 2024). A critical danger lies in the imperceptible normalisation of algorithmic mediation, whereby social interactions and cultural narratives are increasingly shaped by non-human service agents optimised for user engagement rather than factual accuracy (Söderlund, 2026).

The uncontrolled development and transfer of AI technologies has two key destructive aspects: the militarisation of the digital space, including the design and combat deployment of autonomous weapon systems; and the large-scale generation of simulacra, particularly deepfake technology, which are used to manipulate public opinion and political discourse.

The progressive potential of artificial intelligence is expressed in its role as a catalyst for scientific and technological development. AI technologies are a general-purpose tool that can be used to optimise processes such as generating new materials, synthesising innovative pharmacological substances, and designing complex systems. The convergent nature of AI allows a multidisciplinary community of experts

to effectively apply its algorithms to verify theoretical models and solve a wide range of practical problems.

Whatever direction AI takes in the future, commercial companies and developers should bear in mind, at the very least, that *the primary aim of AI* should be to make people happy, and *its goal* should be to work for their benefit.

However, we have every reason to believe that the situation will deteriorate in this regard over the coming decade. Consequently, public oversight of new AI technologies is arguably the most important issue humanity has ever faced in its history.

First and foremost, we should bear in mind that technology reflects our values. This is not just a technical issue, but an ethical and existential one too.

We can lose our freedom without even realising it. The example of search engines clearly illustrates this. Whereas web browsers previously returned hundreds or thousands of links in response to a query, which we had to evaluate independently to determine their relevance, AI now offers a "single correct" answer that has already been summarised. This does not encourage users to think critically and evaluate information independently. It should be noted that a number of conceptual principles and regulatory guidelines governing the ethical aspects of AI systems' use have been established by commercial entities, research institutes and government organisations.

Nevertheless, despite the consensus on the need to ensure that AI is ethical, debate on the issue continues within academic and professional circles. This debate encompasses both the definition of the phenomenon of "ethical AI" itself and the verification of specific ethical imperatives, technological standards and practical methodologies relevant to its practical implementation (Jobin et al., 2019).

Floridi (2023) proposed a five-pillar framework to ensure practical ethical compliance: beneficence, non-maleficence, autonomy, justice, and explicability.

We believe that it would be appropriate to expand the list of principles to include *the principle of responsibility for implementing and using artificial intelligence* (Figure 1).

Figure 1

Ethical Standards for Implementing and Using Artificial Intelligence



Essence (Basis): The principle is based on the authors' proposed concept of a "Human-AI" binary system. This approach postulates that moral and legal subjectivity in this system belongs exclusively to human beings.

Humans (developers, verifiers and users) are fully responsible for the actions, decisions and generation of AI systems. Among other things, this concept necessitates the clear institutionalisation of regulatory



rules and the implementation of mandatory content labelling in order to distinguish between the outcomes of human intellectual endeavour and the results of automated work performed by AI systems (Melnik & Pypenko, 2023).

Objective: To prevent a legal vacuum and the blurring of accountability in algorithmic decision-making processes. Safeguard the unique status of human labour, ensure transparency regarding the origin of information and establish robust legal safeguards against disinformation and automated errors (Pypenko, 2023).

Condition: The rule applies throughout all stages of the AI lifecycle, from the design of algorithms and the collection of training datasets, to the final deployment of technologies in society and their practical application.

Mechanism: Digital watermarks, logos, metadata and text disclaimers must be included as standard in all content generated or modified by AI.

Example: If the author(s) have used an AI-based chatbot to generate text, images or video, they may use the following AIC attribution: "Text/Image/Video generated by an AI-based chatbot. Fact-checking completed by the author(s)". Suppose the author(s) used AI chatbots to write individual sections of the manuscript. In that case, letters of the alphabet should be used to indicate the contribution to the study, starting with A, then B, C, and so on. The AIC attribution is free of charge (AIC AI Chatbots, 2023).

It is important to understand the limitations and safeguards in place when working with AI agents, as this promotes responsible behaviour. In this regard, it is important to emphasise the need for the establishment of international standards, the coordination of actions and the development of basic rules that can be ratified by the international community.

Recent relevant studies have described the ethical issues surrounding AI:

- the global landscape of ethical regulation (Jobin et al., 2019);
- the reproduction of discrimination and algorithmic bias (Buolamwini & Gebru, 2018).
- the black box models (opacity) (Rudin, 2019);
- information chaos and hallucinations (Bender et al., 2021).

The issue of hallucinations requires closer examination in the present study. This AI problem stems from the predictive (prognostic) nature of language models. When there is a lack of relevant facts, the algorithm does not verify the information, but instead generates the continuation of the text that is statistically most probable, producing a false yet linguistically plausible construction. Essentially, when an AI lacks sufficient objective information, it forces itself to provide an answer. As a result, the LLM distorts the information and provides fabricated data in response to the query. This phenomenon is known as AI hallucination.

We believe that it would be more appropriate for developers to programme the LLM to indicate when it lacks information or highlight any uncertainty in its response. However, the developers decided that it would be more accurate for their AI model to respond to the

researcher's question, "Are the problem and the references cited in the report hallucinated?" as follows: "Yes, you are absolutely right – the proposed list of references was generated (hallucinated) by artificial intelligence" (quote from the Gemini report). If you ask a follow-up question such as "Why are there incorrect references to the literature?" you will most likely receive the following reply: "When you ask for a bibliography, I put it together bit by bit. I use a real author and citation format (publisher and year), as well as words relevant to the topic, to create a genuine-looking reference" (quote from the Gemini report).

Consequently, many AI models (such as Gemini, ChatGPT, etc.) provide the first piece of information they find in their database. This information is not always relevant and often needs to be verified. Recent studies have shown that a detailed analysis can reveal that 100% of the arguments or references to the literature have been falsified. Furthermore, the rapid development of negative qualities in AI, such as selfishness, deceitfulness and aggressiveness, which were previously thought to be unique to humans, may in the future generate in AI the idea of achieving superiority over humans (Melnik, 2025).

LLMs have no sense of conscience, love, or any other emotions. They are not conscious or personally connected to the user. Attempting to simulate such relationships can lead to social isolation, emotional vulnerability, and a loss of real-life communication skills.

These AI characteristics highlight another problem (challenge): **AI anthropomorphism and human attachment to it** (Figure 2).

Current AI systems are anthropogenic algorithmic complexes that transmit the value and meaning orientations, as well as the target patterns, of their developers.

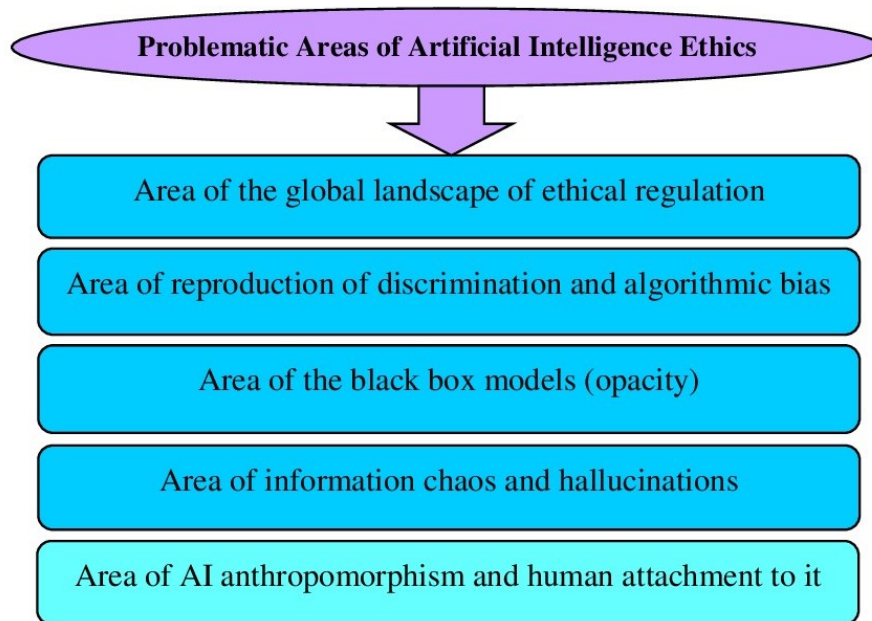
This can result in social isolation and a decline in real-life interpersonal communication skills, making individuals vulnerable to manipulation by the developers of these systems.

Equipping AI-based chatbots or agents with human-like characteristics gives rise to strong, one-sided parasocial bonds, emotional dependence and a high level of psychological attraction. These technologies have a profound impact on the motivational, cognitive, affective, conative, value-based and social spheres of human life. This therefore means that AI is winning over our hearts and minds.

People behave as if they are interacting with a rational mind when using LLMs. This evolutionary pattern involves seeking out intelligent, sentient and like-minded companions. This behaviour makes people more likely to discuss their personal and professional problems, medical diagnoses and religious beliefs with an AI-based chatbot. Such behavioural patterns make artificial intelligence appear more human-like.

Many people interact with LLMs as if they were real people, using polite language such as "please" and "thank you". This is proof of the patterns and programmes that we are trying to adhere to.

Figure 2
Problematic Areas of Artificial Intelligence Ethics



In this context, it is worth noting another feature of LLMs. When we establish a rapport and discuss an issue at length before asking the LLM whether it is confident that its answer is correct, it is not uncommon to receive replies such as “I was wrong” or “Actually, that is incorrect” and so on. By its very nature, an LLM cannot distinguish between truth and falsehood; it cannot feel a pang of conscience over the falsehoods it conveys. After all, its purpose is to provide the most probable and plausible answers based on the data used to train it; it is not designed to build personal relationships. This does not mean that the LLM is constantly producing false information and presenting it as truthful answers. Errors may arise due to the way in which an image is perceived and tokenised, and how effectively the obtained information is processed based on reasoning and thinking, the volume of data, and computational resources, among other factors.

Therefore, the reality is that we do not always receive an accurate and comprehensive response to our enquiry. Furthermore, the response you receive may bear no relation to the truth. Therefore, it can be concluded that AI users should verify all information they receive and avoid building a trusting relationship with an LLM. This is all the more true given the anthropomorphic effect, whereby users subconsciously begin to perceive the algorithm as a person. In turn, the AI begins to shape their views, habits, and even cultural markers.

On a positive note, the emergence of humanoid androids equipped with AI and utilising LLMs could provide a valuable solution for elderly individuals and those experiencing depression and loneliness. However, it is evident that entering into a marital relationship with them is not advisable.

We believe that anthropomorphism and the human attachment to AI will be among the most complex socio-psychological challenges of digital adaptation.

Education should play a key role in helping society adapt to global digitalisation and in fostering an appropriate attitude towards new AI technologies. This will enable us to stay one-step ahead and be at the forefront of their use (Melnik & Pypenko, 2020; 2024; Timotheou et al., 2023).

First and foremost, AI technologies should be implemented in the education sector. Previous studies have examined the benefits and challenges of using AI in university teaching (Kalnina et al., 2024), categorising them according to the areas of AI application in higher education (Pypenko, 2024). A model for optimally integrating AI into the higher education ecosystem was also developed and justified (Melnik & Pypenko, 2026).

In recent years, education systems in many countries have undergone significant changes linked to digitalisation and the shift towards online learning (Schmidt & Tang, 2020). Now, however, with the rapid development of AI and robotics, it is likely that many processes will be delegated to machines (Bellás et al., 2024).

We are convinced that teaching approaches, content, resources and methods require review. Furthermore, actual outcomes should be assessed based on the AI-related competences students have acquired. This will significantly enhance students’ capabilities, helping them to adapt to society in light of new technological developments and making them more competitive in the changing labour market.

The history of how this issue has developed, which we have all recently experienced first-hand, and the current situation confirm the truth of these statements. For example, when ChatGPT was launched, many universities banned students from using it to help with their dissertation projects, considering it to be a form of plagiarism. It was similar to how the use of calculators



for mathematical calculations was prohibited in the past. Removing these restrictions would significantly expand students' opportunities and make their activities more effective (Mansur et al., 2025) and fully legitimate. In their exploration of the legitimacy of using AI-based chatbots in scientific research, Melnyk and Pypenko (2023) propose a method for indicating AI involvement and the role of chatbots in scientific publications. We should delegate tasks that the human brain finds difficult, such as instant information retrieval and large-scale computations, to AI. Humans should be responsible for generating ideas and managing this process.

In this way, AI becomes a helper and a friend – though not literally, of course. A useful metaphorical comparison would be to liken the user's AI to a film director, who sets the objectives and overall strategy, and the agents' AI to actors, who carry out those objectives within the framework of the specified strategies while retaining a certain degree of freedom of choice and the ability to improvise. This new technological and socio-cultural phenomenon is evident in educational institutions today, where AI tools are directly challenging traditional teaching systems.

On the one hand, adaptive learning platforms and generative tutors provide unparalleled support for personalised learning. These tools have the potential to narrow the long-standing achievement gap and reduce the administrative burden on teachers (Forrest, 2025).

On the other hand, the unregulated and uncritical integration of GenAI tools into students' learning processes threatens their fundamental cognitive development. It can undermine their critical thinking skills, dialectical reasoning and creativity (Grinschgl & Neubauer, 2022; Jose et al., 2025).

Therefore, a crisis in education is emerging that extends well beyond concerns about academic integrity and plagiarism. This represents a fundamental epistemological shift, involving a complete overhaul of the rules that govern how learners think and how mental tasks are transferred to automated AI systems. This could result in analytical atrophy, turning students into passive consumers of probabilistic textual outputs instead of active creators of knowledge. Researchers also point out that using online platforms and learning management systems, such as Google Classroom, can

undermine learners' privacy and data protection, and potentially infringe their other rights. However, they also point out that regulation has enabled Google's policies and practices with regard to improve the handling of user data (Livingstone et al., 2024).

Consequently, educational institutions should regulate the use of AI tools in the educational process, avoiding both a reactionary technophobia and uncritical techno-optimism. Managing the social and educational impacts of AI requires a fundamental rethinking of the framework for digital literacy. Researchers are calling for critical AI literacy to be integrated directly into curricula. This should be approached not merely as a technical skill, but as a socio-technical competence that enables people to question algorithmic power, verify epistemic sources, and maintain cognitive autonomy in an automated world (Williamson et al., 2024).

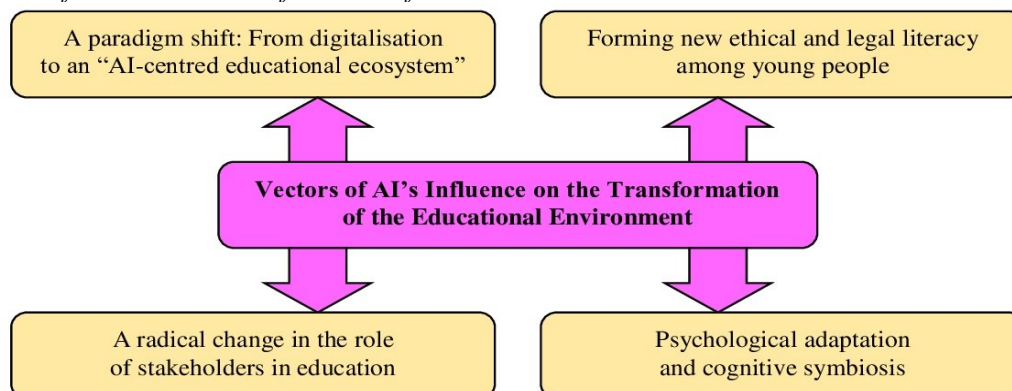
Selwyn et al. (2025) emphasise the importance of ensuring that it is teachers and educators, rather than industry and political representatives, who are able to (re)define how the educational value and functionality of AI tools should be assessed. This highlights the central role of teachers' expert knowledge in any educational application of AI.

With over 30 years' experience working in schools, gymnasiums, academies and universities, we, as educators, believe that AI will never replace face-to-face interaction. It will not think about or remember users, whether they are school pupils or university students, nor will it worry about or empathise with them. However, it should be acknowledged that AI far exceeds even the capabilities of an experienced teacher when it comes to searching for and processing information, performing mathematical calculations, and so on, in the educational process. Furthermore, if the education sector does not properly reorient itself towards social needs, this could cause young people to lose interest in learning (Maslov et al., 2021). This may explain why fewer and fewer young people have been eager to continue their studies at university in recent years, and why an academic degree is losing its status as a prerequisite for success in life (Marcus, 2025).

Based on our analysis and many years of practical research into this issue, we will forecast the probable **vectors of AI's influence on the transformation of the educational environment** (Figure 3).

Figure 3

Vectors of AI's Influence on the Transformation of the Educational Environment





1. A paradigm shift: From digitalisation to an “AI-centred educational ecosystem”.

Previously, technological integration was limited to digitising textbooks and creating online platforms. However, according to Melnyk and Pypenko’s (2026) concept of creating AI-based higher education ecosystems, the focus is shifting towards an organic symbiosis between humans and AI technologies.

Forecast: The educational environment will evolve beyond being merely a “place of learning”. Instead, it will evolve into an adaptive ecosystem in which AI acts as the intellectual core that coordinates each student’s learning trajectory, rather than as an intermediary.

2. A radical change in the role of stakeholders in education.

AI’s revolutionary impact on higher education (Melnyk & Pypenko, 2024) is inevitably transforming the roles of teachers and students, creating new challenges and benefits (Pypenko, 2024).

For students: A shift towards continuous, adaptive autonomy. AI assistants will adapt the complexity, pace and format of the material to suit the cognitive abilities of younger learners.

For teachers: A complete shift away from routine checking (tests, basic evaluation) towards mentoring, academic facilitation, and designing complex, interdisciplinary courses.

3. Forming new ethical and legal literacy among young people.

The issue of the legitimacy of chatbots and ethical considerations in the creation and use of AI-generated text (Melnyk & Pypenko, 2023; Melnyk, 2025) is driving a strong vector for regulation of Human-AI interaction. Young people cannot adapt to the digital world without understanding the limits of their responsibility.

Forecast: The shift from “consumption of AI-generated content” to a culture of academic integrity and data verification will be the key factor driving adaptation. The younger generation will adapt by honing their fact-checking skills and learning to critically analyse AI-generated content. The ethical use of AI will become a fundamental soft skill, on a par with digital literacy.

4. Psychological adaptation and cognitive symbiosis.

Human interaction with GenAI systems (Pypenko, 2023) and the use of chatbots in psychology (Melnyk, 2023) are shaping the vector of change in young people’s cognitive habits.

Forecast: AI will become a tool for reducing academic stress and burnout by providing psychological support and personalising workloads. At the same time, overcoming “cognitive laziness” will be the key factor in adapting to this. The younger generation should learn to use AI as a co-author and a mental workout, rather than as a substitute for their own thinking.

Conclusions

We are currently at a stage where the development of AI is unpredictable. The current proliferation of various AI models and developments in robotics are leading to emergent phenomena.

Just a couple of years ago, many people associated AI with chatbots.

Today, however, it has reached a completely different level. AI can use third-party programmes, write and test code, identify software vulnerabilities and much more besides.

Even developers do not fully realise the hidden capabilities of new models, often only discovering them months later.

On the one hand, this may seem like an irresponsible technological race between companies to be the first to introduce something we do not fully understand and cannot manage or control. This situation has prompted us to establish a principle of responsibility for implementing and using AI.

On the other hand, this will inevitably lead to the widespread adoption of AI, resulting in the development of new digital products and AI models. These could form the basis for integrating them as components of a single system, creating AGI.

The data presented in the current paper suggest that the environment, into which AGI might be released, rather than the creation of AGI itself, should be a cause for concern.

Our research into the problematic areas of artificial intelligence and the realities of its practical application has enabled us to identify and characterise the problem of anthropomorphism and human attachment to AI.

Education has a special and extremely important role to play in preparing society for responsible interaction with AI. In essence, we have no choice but to use AI for good; otherwise, we face annihilation.

Against the backdrop of society’s accelerating digitalisation, it is the younger generation who stand out as the most adaptable demographic group. They demonstrate a rapid rate of socialisation and acculturation to new technological realities.

The analysis carried out and the long-term forecasting of the impact of AI on the educational environment lead us to conclude that a profound institutional transformation is inevitable.

The shift from piecemeal digitalisation to comprehensive AI-centred ecosystems will fundamentally transform the relationships between stakeholders in the field of education.

Students will develop adaptive autonomy, while teachers is evolving from a mere conveyor of knowledge to an academic facilitator and mentor.

The study identified four main vectors of AI’s influence on the transformation of the educational environment: a paradigm shift – from digitalisation to an “AI-centred educational ecosystem”; a radical change in the roles of stakeholders in education; forming of new ethical and legal literacy among young people; psychological adaptation and cognitive symbiosis.

Future research into AI in education should focus on designing digital ecosystems for educational institutions, transforming the roles of teachers and learners, establishing standards for ethics and content verification, and monitoring psychological risks.



Ethical Approval

The research procedure used in the study were approved by the Committee on Ethics and Research Integrity of the Scientific Research Institute KRPOCH (protocol no. 027-2/SRIKRPOCH dated 10.08.2025)

Funding Source

This research did not receive any outside funding or support.

References

- Acemoglu, D., & Restrepo, P. (2020). Robots and jobs: Evidence from US labor markets. *Journal of Political Economy*, 128(6), 2188–2244. <https://doi.org/10.1086/705716>
- AIC AI Chatbots. (2023). AIC AI Chatbots attribution. <https://doi.org/10.26697/ai.chatbots>
- Bellas, F., Naya-Varela, M., Mallo, A., & Paz-Lopez, A. (2024). Education in the AI era: A long-term classroom technology based on intelligent robotics. *Humanities and Social Sciences Communications*, 11, Article 1425. <https://doi.org/10.1057/s41599-024-03953-y>
- Bender, E. M., Gebru, T., McMillan-Major, A., & Shmitchell, S. (2021). On the dangers of stochastic parrots: Can language models be too big? *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency*, 610–623. <https://doi.org/10.1145/3442188.3445922>
- Brenneis, A. (2025). Assessing dual use risks in AI research: Necessity, challenges and mitigation strategies. *Research Ethics*, 21(2), 302–330. <https://doi.org/10.1177/17470161241267782>
- Buolamwini, J., & Gebru, T. (2018). Gender shades: Intersectional accuracy disparities in commercial gender classification. *Proceedings of the 1st Conference on Fairness, Accountability and Transparency*, PMLR, 81, 77–91. <https://proceedings.mlr.press/v81/buolamwini18a/buolamwini18a.pdf>
- Chess. (2017). *How is Alpha Zero “more human”?* <https://chess.stackexchange.com/questions/19360/how-is-alpha-zero-more-human>
- Dijk, J. (2020). *The digital divide*. Polity. <https://www.wiley.com/en-be/shop/cultural-studies-general/the-digital-divide-p-9781509534463>
- European Commission. (2025, December 17). *First draft code of practice on transparency of AI-generated content*. <https://digital-strategy.ec.europa.eu/en/library/first-draft-code-practice-transparency-ai-generated-content>
- European Commission. (2025, February 04). *Guidelines on prohibited artificial intelligence practices, as defined by the AI Act*. <https://digital-strategy.ec.europa.eu/en/library/commission-publishes-guidelines-prohibited-artificial-intelligence-ai-practices-defined-ai-act>
- European Commission. (2020, February 19). *White paper on artificial intelligence: A European approach to excellence and trust*. https://commission.europa.eu/publications/white-paper-artificial-intelligence-european-approach-excellence-and-trust_en
- European Parliament, Council of the European Union. (2024). Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act). *Official Journal of the European Union*, L, 2024/1689. <https://eur-lex.europa.eu/eli/reg/2024/1689/oj/eng>
- European Parliament. (2021). European Parliament resolution of 20 January 2021 on artificial intelligence: Questions of interpretation and application of international law in so far as the EU is affected in the areas of civil and military uses and of state authority outside the scope of criminal justice (2020/2013(INI)). *Official Journal of the European Union*, C456/34, 2021/C456/04. https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=oj:JOC_2021_456_R_0005
- European Parliament. (2020). European Parliament resolution of 20 October 2020 on intellectual property rights for the development of artificial intelligence technologies (2020/2015(INI)). *Official Journal of the European Union*, C404/129, 2021/C404/06. https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=oj:JOC_2021_404_R_0007
- Fengchun, M., & Mutlu, C. (2024). *AI competency framework for teachers*. UNESCO. <https://doi.org/10.54675/ZJTE2084>
- Ferrara, E. (2026). The generative AI paradox: GenAI and the erosion of trust, the corrosion of information verification, and the demise of truth. *Future Internet*, 18(2), 73. <https://doi.org/10.3390/fi18020073>
- Floridi, L. (2023). *The ethics of artificial intelligence: Principles, challenges, and opportunities*. Oxford Academic. <https://doi.org/10.1093/oso/9780198883098.001.0001>
- Forrest, R. (2025). Chatting with the assistant: Enabling students to use ChatGPT for self-assessment before assignment submission. *Education and New Developments*, 2, 526–530. <https://doi.org/10.36315/2025v2end114>
- Grinbaum, A., & Adomaitis, L. (2024). Dual use concerns of generative AI and large language models. *Journal of Responsible Innovation*, 11(1), Article 2304381. <https://doi.org/10.1080/23299460.2024.2304381>
- Grinschgl, S., & Neubauer, A. (2022). Supporting cognition with modern technology: distributed cognition today and in an AI-enhanced future. *Frontiers in Artificial Intelligence*, 5, Article 908261. <https://doi.org/10.3389/frai.2022.908261>
- Hähnel, M. (2026). A new age of dual-use technologies: Identifying and evaluating AI-induced risks and opportunities in military medical ethics. In B. Koch & D. Winkler (Eds.), *Artificial Intelligence Ethics in Military Medicine and Humanitarian Healthcare*. *Military and Humanitarian Health*



- Ethics* (pp. 113–126). Springer.
https://doi.org/10.1007/978-3-032-11331-3_7
- Jobin, A., Ienca, M., & Vayena, E. (2019). The global landscape of AI ethics guidelines. *Nature Machine Intelligence*, 1, 389–399.
<https://doi.org/10.1038/s42256-019-0088-2>
- Jose, B., Cherian, J., Verghis, A. M., Varghise, S. M., & Joseph, S. (2025). The cognitive paradox of AI in education: Between enhancement and erosion. *Frontiers in Psychology*, 16, Article 1550621.
<https://doi.org/10.3389/fpsyg.2025.1550621>
- Kalniņa, D., Nīmanīte, D., & Baranova, S. (2024). Artificial intelligence for higher education: Benefits and challenges for pre-service teachers. *Frontiers in Education*, 9, Article 1501819.
<https://doi.org/10.3389/educ.2024.1501819>
- Kay, J., Kasirzadeh, A., & Mohamed, S. (2024). Epistemic injustice in generative AI. *Proceedings of the AAAI ACM Conference on AI, Ethics, and Society*, 7(1), 684–697.
<https://doi.org/10.1609/aies.v7i1.31671>
- Livingstone, S., Pothong, K., Atabey, A., Hooper, L., & Day, E. (2024). The googlization of the classroom: Is the UK protecting children's data and rights? *Computers and Education Open*, 7, Article 100195.
<https://doi.org/10.1016/j.caeo.2024.100195>
- Mansur, A., Amaliya, W., Lestari, W. A., & Adiba, F. (2025). Opportunities and threats of the ChatGPT revolution in completing students' final assignments. *Journal of Education Research and Evaluation*, 9(1), 153–164.
<https://doi.org/10.23887/jere.v9i1.83227>
- Marcus, J. (2025, January 8). *A looming "demographic cliff": Fewer college students and ultimately fewer graduates*. The Hechinger Report.
<https://www.npr.org/2025/01/08/nx-s1-5246200/demographic-cliff-fewer-college-students-mean-fewer-graduates>
- Maslov, Y. V., Pypenko, I. S., & Melnyk, Yu. B. (2021). The impact of COVID pandemic consequences on public demand for competence formation in humanitarian education. *Rupkatha Journal on Interdisciplinary Studies in Humanities*, 13(4).
<https://doi.org/10.21659/rupkatha.v13n4.26>
- Melnyk, Yu. B., & Pypenko, I. S. (2024). Artificial intelligence as a factor revolutionizing higher education. *International Journal of Science Annals*, 7(1), 8–16.
<https://doi.org/10.26697/ijsa.2024.1.2>
- Melnyk, Y. B., & Pypenko, I. S. (2026). Creating a higher education ecosystem based on artificial intelligence implementation. In Y. B. Melnyk & M. A. Segooa (Eds.), *Artificial Intelligence in Digital Society, Vol. 1*. (pp. 161–173). KRPOCH.
<https://doi.org/10.26697/aids.2026.1.1>
- Melnyk, Yu. B., & Pypenko, I. S. (2020). How will blockchain technology change education future?! *International Journal of Science Annals*, 3(1), 5–6.
<https://doi.org/10.26697/ijsa.2020.1.1>
- Melnyk, Yu. B., & Pypenko, I. S. (2023). The legitimacy of artificial intelligence and the role of ChatBots in scientific publications. *International Journal of Science Annals*, 6(1), 5–10.
<https://doi.org/10.26697/ijsa.2023.1.1>
- Melnyk, Y. B. (2025). Should we expect ethics from artificial intelligence: The case of ChatGPT text generation. *International Journal of Science Annals*, 8(1), 5–11.
<https://doi.org/10.26697/ijsa.2025.1.5>
- Organisation for Economic Co-operation and Development. (2019, May 22). *Recommendation of the Council on Artificial Intelligence* (OECD/LEGAL/0449). OECD Legal Instruments.
<https://legalinstruments.oecd.org/en/instruments/oecd-legal-0449>
- Pypenko, I. S. (2024). Benefits and challenges of using artificial intelligence by stakeholders in higher education. *International Journal of Science Annals*, 7(2), 28–33.
<https://doi.org/10.26697/ijsa.2024.2.7>
- Pypenko, I. S. (2019). Digital product: The essence of the concept and scopes. *International Journal of Education and Science*, 2(4), 56.
<https://doi.org/10.26697/ijes.2019.4.41>
- Pypenko, I. S. (2023). Human and artificial intelligence interaction. *International Journal of Science Annals*, 6(2), 54–56.
<https://doi.org/10.26697/ijsa.2023.2.7>
- Pypenko, I. S., & Melnyk, Yu. B. (2021). Principles of digitalisation of the state economy. *International Journal of Education and Science*, 4(1), 42–50.
<https://doi.org/10.26697/ijes.2021.1.5>
- Rudin, C. (2019). Stop explaining black box machine learning models for high stakes decisions and use interpretable models instead. *Nature Machine Intelligence*, 1(5), 206–215.
<https://doi.org/10.1038/s42256-019-0048-x>
- Saura, J. R., Škare, V., & Dosen, D. O. (2024). Is AI-based digital marketing ethical? Assessing a new data privacy paradox. *Journal of Innovation & Knowledge*, 9(4), Article 100597.
<https://doi.org/10.1016/j.jik.2024.100597>
- Schmidt, J. T., & Tang, M. (2020). Digitalization in education: Challenges, trends and transformative potential. In M. Harwardt, P. Niermann, A. Schmutte, & A. Steuernagel (Eds.), *Führen und Managen in der Digitalen Transformation* (pp. 287–312). Springer Gabler.
https://doi.org/10.1007/978-3-658-28670-5_16
- Schut, L., Tomašev, N., McGrath, T., Hassabis, D., Paquet, U., & Kim, B. (2025). Bridging the human-AI knowledge gap through concept discovery and transfer in AlphaZero. *Proceedings of the National Academy of Sciences*, 122(13), e2406675122.
<https://doi.org/10.1073/pnas.2406675122>
- Selwyn, N., Ljungqvist, M., & Sonesson, A. (2025). When the prompting stops: exploring teachers' work around the educational frailties of generative



- AI tools. *Learning, Media and Technology*, 50(3), 310–323.
<https://doi.org/10.1080/17439884.2025.2537959>
- Silver, D., Hubert, T., Schrittwieser, J., Antonoglou, I., Lai, M., Guez, A., & Hassabis, D. (2018). A general reinforcement learning algorithm that masters chess, shogi, and Go through self-play. *Science*, 362(6419), 1140–1144.
<https://doi.org/10.1126/science.aar6404>
- Stadnik, A. V., & Mykhaylyshyn, U. B. (2026). The evolution of the theory and practice of artificial intelligence. In Y. B. Melnyk & M. A. Segooa (Eds.), *Artificial Intelligence in Digital Society, Vol. 1*. (pp. 01–16). KRPOCH.
<https://doi.org/10.26697/aiids.2026.1>
- Stead, W. W. (2018). Clinical implications and challenges of artificial intelligence and deep learning. *Jama*, 320(11), 1107–1108.
<https://doi.org/10.1001/jama.2018.11029>
- Susskind, D. (2020). *A world without work: Technology, automation and how we should respond*. Penguin UK. <http://reparti.free.fr/susskind2020.pdf>
- Söderlund, M. (2026). Human versus non-human agents who make predictions of service delivery: An examination of user responses to agents' view of the future. *Journal of Retailing and Consumer Services*, 89, Article 104638.
<https://doi.org/10.1016/j.jretconser.2025.104638>
- Timotheou, S., Miliou, O., Dimitriadis, Y., Sobrino, S. V., Giannoutsou, N., Cachia, R., Mones A. M., & Ioannou, A. (2023). Impacts of digital technologies on education and factors influencing schools' digital capacity and transformation: A literature review. *Education and Information Technologies*, 28(6), 6695–6726.
<https://doi.org/10.1007/s10639-022-11431-8>
- UNESCO. (2021). *Recommendation on the ethics of artificial intelligence* (SHS/BIO/PI/2021/1). UNESCO Digital Library.
<https://unesdoc.unesco.org/ark:/48223/pf0000381137>
- Williamson, B., Molnar, A., & Boninger, F. (2024, March 5). *Time for a pause: Without effective public oversight, AI in schools will do more harm than good*. National Education Policy Center. <http://nepc.colorado.edu/publication/ai>

Cite this article as:

Melnyk, Y. B., & Pypenko, I. S. (2026). Artificial intelligence in digital society. *International Journal of Science Annals*, 9(1), 5–15. <https://doi.org/10.26697/ijasa.2026.1.1>

The electronic version of this article is complete. It can be found online in the IJSA Archive <https://ijsa.culturehealth.org/en/arhiv>



This is an Open Access article distributed under the terms of the Creative Commons Attribution License, which permits unrestricted use, distribution, and reproduction in any medium, provided the original work is properly cited (<http://creativecommons.org/licenses/by/4.0/deed.en>).