

ChatGPT as a Source of Information in Higher Education: How are the Results Provided by Generative AI Valued?

ChatGPT como fuente de información en la Educación Superior: ¿cómo se valoran los resultados que proporciona la IA generativa?

Lucila Finkel; Pilar Parra Contreras; Yolanda Martínez Solana; Carla Matos Mejías

Recommended citation:

Finkel, Lucila; Parra-Contreras, Pilar; Martínez-Solana, Yolanda; Matos-Mejías, Carla (2024). "ChatGPT como fuente de información en la Educación Superior: valoración de los resultados que proporciona la IA generativa". *Profesional de la información*, v. 33, n. 6, e330604.
<https://doi.org/10.3145/epi.2024.0604>

Artículo recibido el 08-09-2024
Aceptación definitiva: 13-12-2024



Lucila Finkel ✉

<http://orcid.org/0000-0002-3742-0602>
Universidad Complutense de Madrid
Instituto Complutense de Sociología para el
Estudio de las Transformaciones Sociales
Contemporáneas (TRANSOC)
Campus de Somosaguas, s/n
28223 Pozuelo de Alarcón (Madrid), España
lfinkel@ucm.es



Pilar Parra-Contreras

<http://orcid.org/0000-0001-6119-5110>
Universidad Complutense de Madrid
Instituto Complutense de Sociología para el
Estudio de las Transformaciones Sociales
Contemporáneas (TRANSOC)
Campus de Somosaguas, s/n
28223 Pozuelo de Alarcón (Madrid), España
pparra@ucm.es



Yolanda Martínez-Solana

<http://orcid.org/0000-0001-9164-4634>
Universidad Complutense de Madrid
Facultad de Ciencias de la Información
Av. Complutense, 3
28040 Madrid, España
mymartinez@ccinf.ucm.es



Carla Matos-Mejías

<http://orcid.org/0000-0003-0339-4351>
Universidad Complutense de Madrid
Facultad de Ciencias de la Información
Av. Complutense, 3
28040 Madrid, España
carmatos@ucm.es

Resumen

La irrupción de ChatGPT 3 en noviembre de 2022, y sus distintas versiones posteriores hasta la actual 4-o, y su expansión y creciente utilización masiva, está planteando diferentes cuestiones de especial interés relacionadas con su uso como fuente de información en la Educación Superior. Siendo conscientes de que estamos asistiendo a la generalización de la Inteligencia Artificial (IA) generativa, este artículo parte de la idea de que es fundamental comprender cómo los estudiantes universitarios percibían y evaluaban ChatGPT 3 a los pocos meses de su lanzamiento. La investigación, llevada a cabo con una metodología mixta, analiza con métodos bivariados y de segmentación múltiple las 4.800 respuestas de estudiantes universitarios a una encuesta *online*, en la que debían indicar su grado de conocimiento y experiencia previa con la IA y sus valoraciones sobre la información proporcionada por ChatGPT. Por otro lado, a partir de un ejemplo sobre los atentados del 11-M, se ha realizado un análisis temático de las respuestas de ChatGPT, con una valoración experta. Los resultados obtenidos muestran que existen diferencias significativas en el grado de información del alumnado y su experiencia previa con ChatGPT, así como en la valoración sobre la información que aporta ChatGPT, en relación a algunas variables sociodemográficas y educativas. Se constata asimismo que la experiencia previa de uso con la IA generativa incide en su valoración como fuente de información o como herramienta de difusión de bulos o *fake news*. Por último, el contraste de sus puntuaciones con la valoración experta pone de relieve la necesidad de enseñar a los estudiantes a entrenar adecuadamente ChatGPT para obtener mejores respuestas y valorarlas críticamente. Esta aproximación puede servir de base para el diseño de estrategias formativas que permitan introducir la IA generativa en la Educación Superior.



Palabras clave

ChatGPT; Inteligencia Artificial; IA; Información; Educación superior; Estudiantes universitarios

Abstract

The release of ChatGPT-3 in November 2022, followed by its subsequent versions up to the current 4.0, and its rapid expansion and widespread use, has raised several questions of particular interest regarding its role as an information source in Higher Education. Recognising the growing generalisation of generative Artificial Intelligence (AI), this article is based on the premise that it is essential to understand how university students perceived and evaluated ChatGPT-3 a few months after its launch. The research conducted using a mixed-methods approach, analyses 4,800 university student responses to an online survey using bivariate and multiple segmentation methods. The survey asked participants to state their level of knowledge and experience with AI, as well as their evaluations of the information provided by ChatGPT. Additionally, experts conducted a thematic analysis of ChatGPT's responses using an example based on the 11-M attacks. The findings reveal significant differences in students' levels of knowledge, prior experience with ChatGPT, and their evaluations of the information it provides, linked to various sociodemographic and educational variables. The results also underline that prior experience with generative AI influences its evaluation, both as an information source and as a tool for disseminating fake news. Finally, the comparison of student evaluations with those of experts stresses the need to teach students how to effectively train ChatGPT to generate better responses and how to critically evaluate them. This approach could serve as a foundation for designing educational strategies to integrate generative AI into Higher Education.

Keywords

GPT Chat, Artificial Intelligence; AI; Information; Higher Education; University Students.

1. Introduction

In the digital age, technology has drastically transformed the access to and consumption of information (**Micó et al.**, 2022). The advent of Artificial Intelligence (AI) language models has revolutionised the way people access information and media (**Grizzle et al.**, 2023; **Ocaña-Fernández et al.**, 2019; **Túñez-López et al.**, 2021). Among the tools capable of engaging in text-based conversations using Artificial Intelligence, ChatGPT (an acronym for Generative Pre-trained Transformer) developed by OpenAI, has revolutionised the sector. Following the release of version 3 in November 2022, it reached one million registered users within just five days. According to a report published by the Australian newspaper *Financial Review* (**Sier**, 2022), Netflix took 40 months to reach that figure, Facebook 10 months, Twitter 24 months, and Instagram three months. In March 2023, OpenAI launched the GPT-4 model, an improved version that required a paid subscription, and in May 2024, it introduced an update, GPT-4.0, which was more powerful and featured greater capabilities than its predecessors.

ChatGPT is an example of a generative Artificial Intelligence language model, or LLM, developed by OpenAI (**Márquez**, 2023). It is based on a Transformer neural network architecture (**Vaswani et al.**, 2017), which has proven effective in generating natural language across a variety of tasks, including making machine translations and generating dialogue. ChatGPT's neural network is trained using large text datasets, enabling it to infer semantic relationships between words and learn natural language patterns. As a result, it can generate coherent and contextually relevant text based on the input provided by the user (**Pavlik**, 2023).

Since its launch, the use of ChatGPT has become widespread across various fields. However, its use as a substitute for educational information sources or even traditional media presents important challenges and potential consequences for how people learn, interact, and make decisions (**Flores-Viva; García-Peñalvo**, 2023). This is because the responses are generated by a machine rather than a person with knowledge and expertise in the subject, who can assess the impact and validity of the information. Additionally, it has been highlighted that such AI tools produce persuasive and coherent texts, which can influence public opinion—understood as the set of attitudes, beliefs, and values held by members of a society about a specific topic or issue (**Bennett et al.**, 2006). For this reason, some authors advocate using these tools as a complement to other information sources rather than as a replacement for human expertise and specialised knowledge (**Li et al.**, 2021).

Being aware that we are witnessing the early stages of what will, in the coming years, become the true expansion and widespread adoption of AI, and in contrast to other studies focused on the perceptions of the general population (**Centeno-Martín et al.**, 2023), the opinions of academics regarding ChatGPT (**Bringula**, 2023), or those of communication professionals (**Gutiérrez-Caneda et al.**, 2023; **Micó et al.**, 2022; **Sánchez-García et al.**, 2023), it is especially interesting to explore how young university students are responding to generative AI. It is also crucial to explore the advantages and risks of using ChatGPT for information retrieval and to assess the quality of the responses it provides, given its likely and increasing use as an assistant in education and for future professionals in the field of communication.

2. Factors Related to ChatGPT's Search System and Information Generation

Artificial Intelligence is understood as the interdisciplinary field of science and engineering focused on the design and development of algorithms and computer systems capable of performing tasks for which human beings require intelligence (Russell; Norvig, 2021). This can include tasks such as recognising patterns, processing natural language, decision making, and solving complex problems. As these models become more advanced, they can generate responses that are increasingly natural and difficult to distinguish from those generated by humans. This may mean that users can ask questions and receive relevant and coherent answers without having to look for the information elsewhere. (García, 2023). To obtain the most effective responses from ChatGPT, it is advisable to use prompts, which are a set of techniques for engaging effectively with large language models (LLMs) like ChatGPT. Prompts involve specific tags—clear instructions or keywords—that help steer the model's responses toward particular questions or tasks (Giray, 2023; White *et al.*, 2023). However, the accuracy and objectivity of the responses generated by AI language models depend on various factors. The most significant is that the model has been trained on a high-quality dataset that encompasses a wide range of knowledge and perspectives on the topic (Lample; Conneau, 2020).

Although Artificial Intelligence language models trained on large amounts of text generate convincing responses that can influence people's attitudes and behaviours, it is important to note that they lack the ability to understand the context or intent behind a question or comment (Radford *et al.*, 2019). The lack of experience and critical judgement that human beings can contribute, besides the absence of a deep understanding of the social and political context can produce incorrect or incomplete answers, especially where the context is complex or ambiguous. In other cases, the model may be trained on an unrepresentative, or biased data set, so the model may not generate accurate or unbiased answers.

In recent literature, the production of misinformation has emerged as one of the most extensively discussed issues. Among other concerns, it has been stressed that AI does not rely on official sources and does not include references to identify the origin of its information unless explicitly requested to do so. As Polo (2023) points out, "one of the main challenges faced by this model is the lack of transparency regarding the source and the accuracy of the information provided." In the academic world, AI also has two fundamental limitations: firstly, it presents incorrect answers as true, and secondly, it is unable to produce academic work with properly applied and reliable citation systems. On the other hand, by using content generated and available on the internet, the information it provides may be true, false, or inaccurate (García, 2023). Furthermore, the tool may occasionally fabricate information, including non-existent scientific references (Walters; Wilder, 2023). ChatGPT's dialogue format enables it to acknowledge its mistakes and even apologise for them.

Misinformation is often also associated with issues of bias, as the vast datasets used to train ChatGPT may include biased information that leads to skewed responses (Rodríguez, 2023). Since there is no evaluation of truth or falsehood, AI programs can perpetuate historical biases or inaccuracies present in their training data. According to a UNESCO report, some of these biases become particularly significant when artificial language models are trained on data that reflect systemic inequalities or discrimination. In addition to prejudices and stereotypes, the model may also perpetuate racist, sexist, homophobic, ableist, and other discriminatory language (Sabzalieva; Valentini, 2023). Another characteristic of ChatGPT's responses is the absence of instant verification of the information it provides, as it does not perform web searches. Furthermore, in its initial version 3, the information provided only extended up to 2021, meaning that the information offered by ChatGPT on many topics was not up to date.

These issues, among others, have sparked growing interest in understanding users' perceptions and experiences regarding the use of ChatGPT. An analysis of the most relevant publications identified in 2023 in the Scopus and Google Scholar databases, using keywords such as "ChatGPT," "education", and similar terms, reveals that one of the prominent categories is user motivation and experiences, accounting for 10.9% of the publications (Saz-Pérez; Pizà-Mir, 2024). Regarding the analysis of university students' perceptions and experiences, there are several noteworthy reviews of the state of AI in the educational field (García-Peñalvo *et al.*, 2024; Yusuf *et al.*, 2024) and research that explores the potential and consequences of AI and ChatGPT in Higher Education in Spain (García-Peñalvo, 2023; Rodríguez-Rivera *et al.*, 2023; Segovia-García, 2023), as well as in other countries such as Norway (Grassini, 2023), Turkey (Sart, 2023), and Bulgaria (Valova *et al.*, 2024), and on other continents, including Mexico (García-Sánchez, 2023), India (Das; Madhusudan, 2024), the Philippines (Obenza *et al.*, 2024), and Vietnam (Ngo, 2023). These studies share a common focus on analysing how university students perceive the use of ChatGPT for learning, including its benefits, barriers, and potential solutions.

Studies have also been conducted to test ChatGPT and its reliability, primarily in areas such as research, education, medicine, and communication (Morath *et al.*, 2024; Pavlik, 2023; Stokel-Walker; Van Noorden, 2023), with the aim of evaluating the accuracy and consistency of its responses. It is reasonable to assume that this line of research will lead to numerous publications in other fields in the coming years. The objective of this article is to explore university students' perceptions and experiences with ChatGPT during its initial stages through an online survey and to examine the validity of ChatGPT's responses using an example analysis that can be utilised and replicated in university education. To this end, the analysis is structured into two main sections, corresponding to two distinct methodological approaches: the first presents

the survey results, based on the primary hypothesis that university students' perceptions of ChatGPT depend on their level of knowledge about it, prior usage experience, as well as their age and the academic discipline of their degree programme.

Studies have also been conducted to test ChatGPT and its reliability, particularly in fields such as research, education, medicine, and communication (Morath *et al.*, 2024; Pavlik, 2023; Stokel-Walker; Van Noorden, 2023). The purpose of these studies is to evaluate the precision and consistency of its responses. It is to be expected that this line of research will result in numerous publications in various fields in the years to come. The second section analyses ChatGPT's responses by comparing answers to generic questions that might be asked by a university student with no prior knowledge of generative AI, against responses to questions formulated with prompts, following recommendations from experts in the university field (Mollick; Mollick, 2023; Torres-Salinas; Arroyo-Machado, 2023; UC3M, 2023). The analysis uses as an example a landmark event in Spain's history: the attacks of 11 March 2004, a topic studied in depth by Professor Martínez Solana, the author of this article. The comparison between students' evaluations and expert evaluations—conducted on a specific topic and with an understanding of how to use generative AI—highlights the importance and necessity of reflecting on the informational use of ChatGPT to achieve its maximum potential.

3. Methodology

The research, conducted at a strategic moment shortly after the launch of ChatGPT, employed a mixed-methods design. Students' perceptions and evaluations were explored using an online survey (Arroyo Menéndez; Finkel Morgenstern, 2019), while the accuracy of ChatGPT's responses was assessed through qualitative thematic analysis. The online survey, implemented using the Limesurvey application, was distributed via email to all students enrolled in official and proprietary degree programmes at the Complutense University of Madrid during May and June 2023. Because the survey was accessible on various devices (tablets, laptops, and mobile phones), incentives were offered—specifically, three wireless headphones raffled among participants who completed the questionnaire (Díaz-de-Rada, 2022)—and because the student population showed significant interest in Artificial Intelligence, particularly ChatGPT, a high response rate was achieved. 7,196 responses were received (4,800 of which were complete), representing a relatively high response rate of 11.18% for this type of study, given a total student population of 64,385 during the 2022–23 academic year (Sánchez-Carrión *et al.*, 2012). Since the sampling was probabilistic, with data on the total population and its distribution by type of study and faculty provided by the Integrated Institutional Intelligence System (SIDI), partial response rates, sampling errors, and weighting factors were calculated to ensure the representativeness of the sample obtained. It is important to note that this study may have been affected by voluntary response bias, also known as self-selection bias, where respondents with specific characteristics are more likely to participate in the research than others.

The approach to understanding students' perceptions was primarily based on the averages of various semantic differential scales (Osgood *et al.*, 1957), which were used in this case to measure the connotative or subjective meaning that participants attribute to the information provided by ChatGPT. Based on the literature review, the semantic pairs selected to evaluate this information are as follows: Incomplete/Complete, Doubtful/Reliable, Biased/Unbiased, Obsolete/Updated, and Incorrect/Correct. These five scales were measured on a range from 0 to 10. To analyse the survey data, conducted using SPSS 28, the main descriptive statistics of the variables considered were identified, and bivariate analyses were performed, including statistical tests for associations between variables (χ^2) and tests for differences in means for variables with three or more response categories (one-way ANOVA). Additionally, the multivariate segmentation analysis technique (Chaid) was applied to identify student profiles based on the evaluation of their experience with ChatGPT.

As previously mentioned, the thematic analysis of ChatGPT's responses was conducted using the 11 March 2004, attack in Madrid, commonly known as "11-M," as the subject. This topic was chosen because one of the researchers involved in this project is a specialist in the field and the author of the article "*La credibilidad de la fuente oficial. La secuencia 11-14 M*" (*The Credibility of the Official Source. The sequence "11–14 M"*) (Martínez-Solana, 2005) and co-writer of the book "*Los atentados del 11-M. Cuatro días que cambiaron la historia de España*" (*The 11-M Attacks: Four Days That Changed the History of Spain*) by Martínez-Solana and Manchón-Campillo (2022) provide a detailed documentary analysis, including official archives and relevant journalistic articles, serving as expert information on the topic. Additionally, 2024 marks 20 years since the 11-M attacks in Madrid, which left 192 dead and nearly 2,000 injured, profoundly impacting Spanish society, particularly regarding the media coverage, which became known as the "11-M conspiracy theory." University students today, who did not experience the events firsthand, can only learn about this topic through information sources, including ChatGPT.

Seven questions were designed from two perspectives: one from the point of view of a student without knowledge of how to use generative AI asking a straightforward and simple question, and the other from the perspective of someone with basic training in using ChatGPT. In the latter case, prompts were crafted using various strategies, such as requesting the simulation of different roles or specific formats, performing specific tasks like summarising the topic, citing information, using a particular type of language or literary style, drafting a news article, or even asking ChatGPT itself to create a well-constructed prompt. Based on ChatGPT's responses, an expert thematic analysis of the content was conducted (with and without prompts). Each response was evaluated on a scale from 0 to 10 using the same semantic

differential scales employed in the earlier quantitative section: Incomplete/Complete, Doubtful/Reliable, Biased/Unbiased, Obsolete/Updated, and Incorrect/Correct.

4. Perception of the Degree of Knowledge and Assessment of the Experience of Using ChatGPT

Following an exhaustive screening process that reduced the number of cases to those who had completed the questionnaire in its entirety ($n=4,800$), the sample was cross-weighted by faculty and level of study (degree, double-degree, master's and doctorate). One of the first questions in the questionnaire, referring to whether the respondent had heard of any of the different AI applications considered, acted as a filter question, since if the respondent stated that they had not heard of ChatGPT, the questionnaire ended. In this way, 96.7% of the sample acknowledges having heard of various voice recognition technologies (such as Siri, Alexa, Google Home, etc.) and 95% of ChatGPT. This high level of knowledge is followed by other tools such as the translation application DeepL with 40.3%, followed by the image creation tools Dalle2 and Midjourney, with 24.3% and 23.1% respectively.

The question regarding the frequency of use of these applications shows, however, that the only applications that are used by around 50% of the sample are voice recognition technologies (Siri, Alexa, Google Home, etc.) and ChatGPT. In the case of the latter, when we analyse the frequency of use, it is found that 53.4% of the sample has never or almost never used it, compared to 27.4% who have used it sometimes and 19.2% who use it frequently and very frequently, truly surprising percentages if one takes into account that ChatGPT was launched at the end of November 2022, and at the time the fieldwork for the survey was carried out (May and June 2023), barely six months had passed. To analyse the survey data, it is necessary to first characterise the group of students who report using ChatGPT with varying degrees of intensity. Among the sociodemographic variables associated with its use, those showing statistical significance ($p<0.05$) include gender, with a higher proportion of men using it (57.3%), age, with young adults aged 21–25 being the most frequent users (while those aged 18–20 use it the least), and social class.

Regarding this last variable, a clear distinction is observed among social classes in terms of ChatGPT usage. The working class shows the lowest usage, with only 44.3% of students using ChatGPT compared to 55.7% who do not. Similarly, usage among the lower-middle and middle classes is 46.3% and 45.1%, respectively. In contrast, usage increases to 50.8% for the upper-middle class and 59.6% for the upper class. Notably, there is a 15.3 percentage-point difference between the working class and the upper class. Academic variables also show statistical significance ($p<0.05$) in relation to the use of ChatGPT. Among fields of study, students in Engineering and Architecture stand out, with 72% using ChatGPT compared to 27.8% who do not. They are followed by students in Sciences and Social and Legal Sciences, with usage rates of 53% and 50.1%, respectively, while students in Arts and Humanities and Health Sciences show lower usage rates of 38.3% and 37.6%.

By faculty, the highest percentages of ChatGPT usage—above 50%—are found in Computer Science (81.1%), Statistical Studies (73%), Commerce and Tourism (58.6%), Mathematical Sciences (58.1%), Physical Sciences (57.9%), Information Sciences (55.0%), and Biological Sciences (50.3%). Regarding the type of studies, master's students are the highest users of ChatGPT at 52.6%, followed by students in double-degree programmes, doctoral programmes, and bachelor's degrees. On the other hand, among students who have heard of ChatGPT but have not used it, the category with the highest percentage is “it doesn't interest me,” at 37.8%. This result aligns with other studies that also show a high level of disinterest among the adult population (Centeno-Martín *et al.*, 2023). It is important to note that these studies, both among the adult population and university students, are among the first survey-based investigations being conducted in Spain. As such, they represent early trends that will require further exploration in the coming months.

The second most notable response is a lack of trust in ChatGPT, with 24.6% selecting the option “I don't think it's reliable.” An analysis by the field of study of the students' degree programmes reveals that the highest levels of disinterest are found among students in Sciences (40%) and Arts and Humanities (39.6%). Meanwhile, for the item “I don't think it's reliable,” the highest percentages are recorded among students in Engineering and Architecture (38.6%), whereas students in Health Sciences and Sciences show the least distrust, with 18.5% and 20.9%, respectively. The analysis of AI usage requires, in addition to characterising its users, a deeper exploration of their perceptions and experiences with ChatGPT. These are reflected in the indicators, questions, and descriptive statistics in Table 1, which show that the average self-perceived level of knowledge (5.57) is significantly lower than the evaluation of the experience itself (7.28).

Table 1: Descriptors of Self-perception of the Degree of Knowledge and Assessment of the user Experience of ChatGPT.

Indicator	Question	Mean	Std. dev.	Valid n
Self-perception of the degree of knowledge of ChatGPT	On a scale of 0 to 10, with 0 being "I have no idea" and 10 "I am very informed", do you think you really know how ChatGPT works?	5.57	2.43	2.117
Assessment of the user experience of ChatGPT	Overall, how do you rate your experience with ChatGPT being 0 very negative and 10 very positive?	7.28	1.77	2.116

Source: Own elaboration.

A detailed analysis of the averages for both indicators in relation to gender, age, and field of study shows that the averages differ

across the various categories considered (Table 2). Men rate their knowledge and experience with ChatGPT higher than women, while individuals identifying as non-binary fall in an intermediate position between the two. Regarding age, the youngest students and those over 26 report higher levels of self-perceived knowledge, though the latter group is slightly less enthusiastic about their experience compared to those under 22. In terms of field of study, students in Engineering and Architecture stand out significantly above the rest, both in their self-perception of knowledge and in their evaluation of their experience.

Table 2: Averages of the Self-perception of Knowledge Scale and Evaluation of ChatGPT usage Experience by Gender, Age, and Field of Study of the Degree Programme (*).

		Self-perception of the Degree of Knowledge of ChatGPT	Assessment of the user Experience of ChatGPT
Sex / Gender	Male	6.18	7.40
	Female	5.07	7.20
	Non-binary	5.30	6.77
Age	18-20 years	5.60	7.41
	21-22 years old	5.50	7.42
	23-25 years	5.35	7.26
	26 years or older	5.80	7.01
Field of knowledge	Arts and Humanities	5.54	7.04
	Social and Legal Sciences	5.61	7.26
	Science	5.60	7.28
	Health Sciences	5.05	7.41
	Engineering and Architecture	6.51	7.49

Source: Own elaboration.
 (*) Anova test for the difference of means $p < 0.05$ in all variables considered.

These initial evaluations require further exploration. To address this, the questionnaire also included a question allowing students to assess the information provided by ChatGPT using a set of items based on Osgood's semantic differential scale. Table 3 shows that the average scores range from 5.71 on the "biased – unbiased" scale to 6.62 on the "incomplete – complete" scale. This suggests a fairly balanced perception between the opposing terms, with a slight lean (above 5) towards the positive end of the scale in all cases.

Table 3: Descriptives of the Evaluation of ChatGPT Information.

Question	Semantic Differential	Mean	Std. dev.	Valid n
In your experience and using a scale from 0 to 10, how would you rate the information that ChatGPT has provided to you?	Incomplete (0) – Complete (10)	6.62	2.20	2.114
	Doubtful (0) – Reliable (10)	5.75	2.33	2.088
	Biased (0) – Unbiased (10)	5.71	2.37	2.071
	Obsolete (0) – Updated (10)	6.44	2.48	2.082
	Incorrect (0) – Correct (10)	6.01	2.34	2.087

Source: Own elaboration.

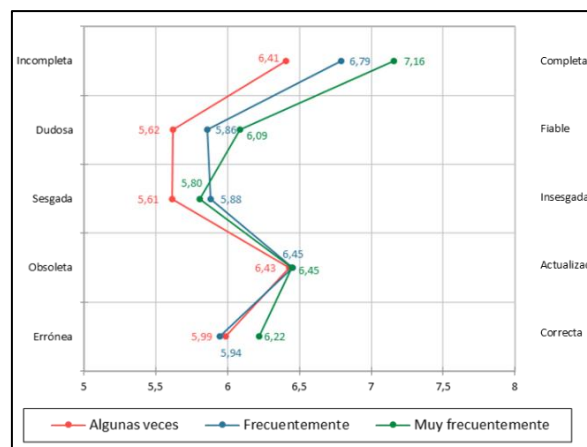


Chart 1: Average Ratings of the Information Provided by ChatGPT According to its Frequency of Use (*).

Source: Own elaboration.

(*) ANOVA test for the difference of means $p < 0.05$ in the incomplete-complete, doubtful-reliable and biased-unbiased scales.

However, it is important to analyse the average evaluations of the information provided by ChatGPT based on the frequency of its use (considering the categories "sometimes," "frequently," and "very frequently," while omitting the responses "never" and "almost never"). Graph 1 illustrates a more positive evaluation across all scales as the frequency of use increases. Those who use ChatGPT more frequently perceive the information as tending to be more complete (average of 7.16), more reliable (6.0), less biased (5.8), more up to date (6.4), and more accurate (6.2). It is also important to identify how evaluations of ChatGPT vary according to the sociodemographic or academic variables

considered. As shown in Graph 2, as age increases, the ratings on the five scales considered progressively decrease. Regarding social class, particularly high ratings are observed among students who self-identify as belonging to the upper class, with ratings gradually decreasing as one moves down the social class scale.

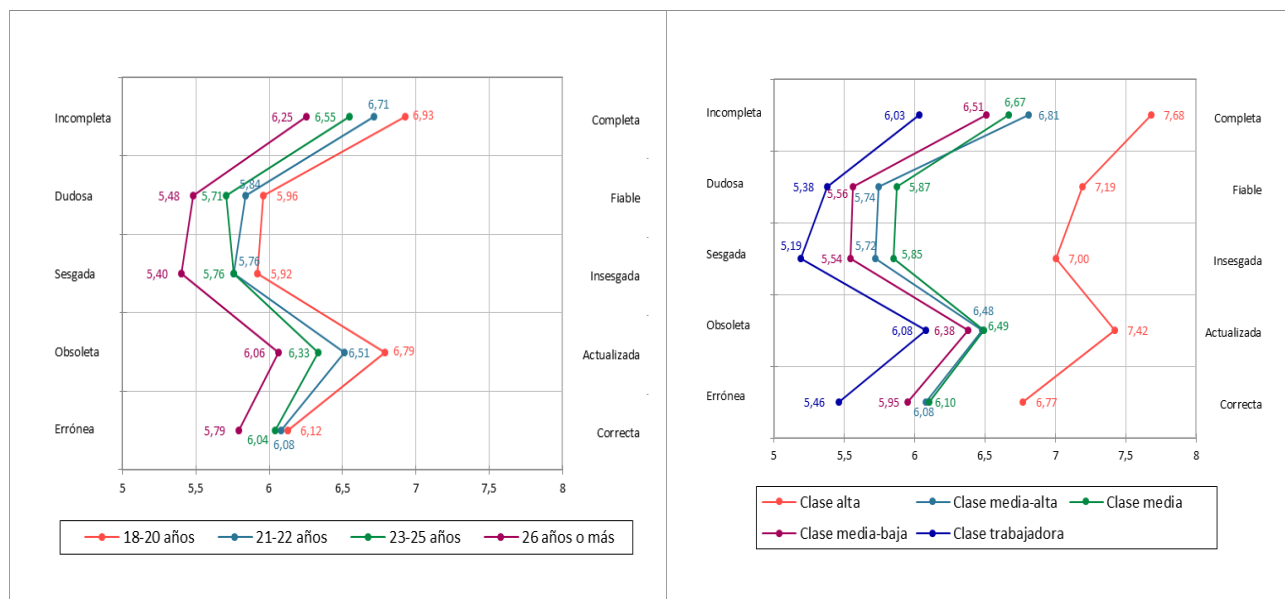


Chart 2: Average Ratings of the Information Provided by ChatGPT According to Age and Social Class (*)

Source: Own elaboration.

(*) ANOVA test for the difference of means between age groups: $p < 0.05$ on the incomplete-complete, doubtful-reliable biased-unbiased and scales. Between social classes: $p < 0.05$ on all scales.

In the case of academic programmes (Chart 3), aligning with the higher ratings from younger students observed in Table 2, it is evident that students pursuing bachelor's and double-degree programmes score higher on all scales compared to those enrolled in master's and doctoral programmes. Regarding the field of study, students in Health Sciences and Sciences provide the highest ratings, while those studying Arts and Humanities or Social and Legal Sciences give the lowest evaluations.

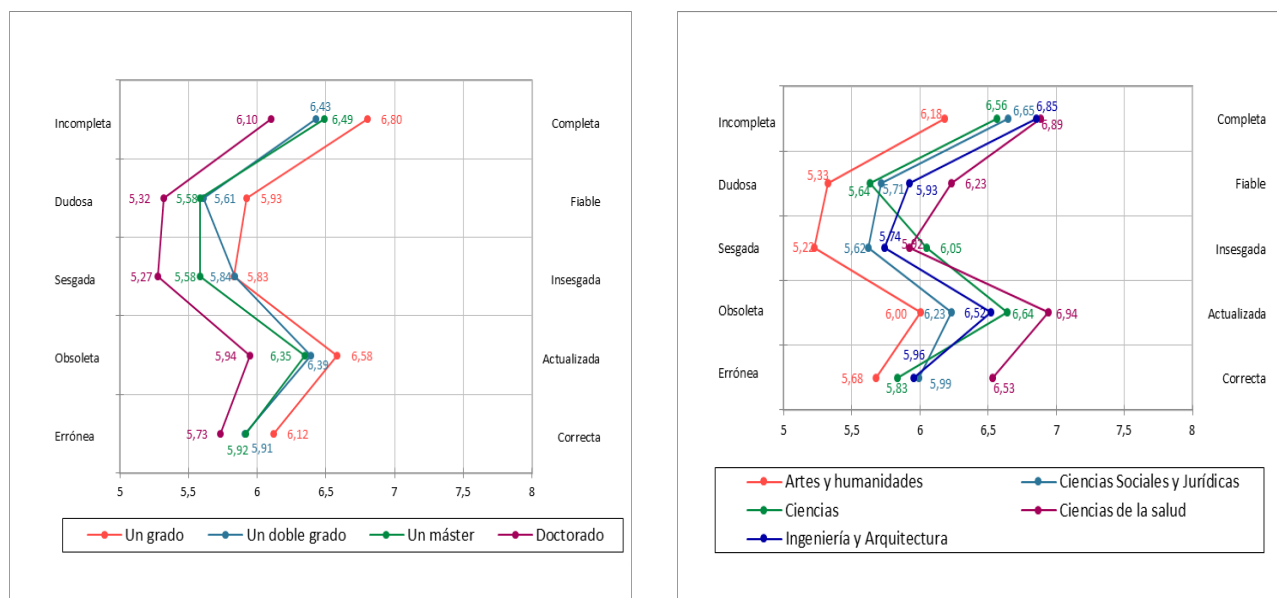


Chart 3: Average Ratings of the Information Provided by ChatGPT by Type of Study and Field of Knowledge of the Degree (*).

Source: Own elaboration.

(*) ANOVA test for the difference of means between the different studies: $p < 0.05$ for the scales incomplete-complete, doubtful-reliable, biased-unbiased, and obsolete-updated. Between the different branches of knowledge: $p < 0.05$ on all scales.

It is important to correlate the average ratings of ChatGPT with the two dependent variables considered earlier: the level of knowledge about its functionality and the evaluation of prior experience. Both variables have been recoded into three categories to facilitate analysis. Tables 4 and Table 5 reveal that both variables play a significant role in the assessment of various aspects related to the program's responses. It can be argued, therefore (Table 4), that greater

knowledge about generative AI, and more notably, greater exposure to or prior experience with it (Table 5), increase the likelihood of having a more positive perception of ChatGPT across all scales considered.

Table 4: Evaluation of the Information Provided by ChatGPT based on the Degree of Information on its Functioning.

	Not at all or Very Poorly Informed	Moderately Informed	Quite or Very Informed
Incomplete Complete	6,12	6.60	6.93
Doubtful Reliable	5.22	5.80	6,00
Biased Unbiased	5.40	5.79	5.80
Obsolete Updated	6.26	6.46	6.51
Incorrect Correct	5.62	6.05	6.19
Source: Own elaboration.			
(*) ANOVA test for the difference of means $p < 0.05$			

Table 5: Evaluation of the Information Provided by ChatGPT based on the Assessment of the Experience of Use.

	Bad or Very Bad	Intermediate	Good or Very Good
Incomplete Complete	2.91	5.47	7.17
Doubtful Reliable	2.40	4.69	6.26
Biased Unbiased	3.33	5.03	6.04
Obsolete Updated	4.45	5.74	6.75
Incorrect Correct	2.42	4.92	6.53
Source: Own elaboration.			
(*) ANOVA test for the difference of means $p < 0.05$			

The analysis of students' initial use and evaluation of ChatGPT is further enriched by examining their level of agreement or disagreement with specific statements regarding the informational use of ChatGPT, such as its ease of accessing information or its potential issue of contributing to the spread of fake news. Table 6 highlights a significant relationship between these variables and the level of knowledge about ChatGPT's functionality. A positive association (Gamma statistic, γ) is observed between the level of knowledge and the perception of ease of staying informed, while a negative association is observed between the level of knowledge and the perceived risk of spreading fake news. These associations are also evident in the evaluation of ChatGPT usage experience (Table 7), reinforcing the notion that knowledge and exposure contribute to more favourable assessments of various aspects of generative AI.

Table 6: Level of Agreement on Informational Usefulness based on the Level of Knowledge on its Functionality.

Question		Not or Very Little Informed	Averagely Informed	Quite or Very Informed
With this type of tool, it is easier to be informed (*)	Disagree	25,2%	19,3%	18,3%
	Indifferent	25,8%	24,8%	20,4%
	Agreed	49,0%	56,0%	61,3%
	TOTAL	100,0%	100,0%	100,0%
This type of tool encourages the spread of hoaxes or fake news (**)	Disagree	153,4%	17,5%	24,2%
	Indifferent	27,7%	11,7%	27,0%
	Agreed	56,8%	52,9%	48,8%
	TOTAL	100,0%	100,0%	100,0%
Source: Own elaboration.				
(*) Significant association (χ^2), $p < 0.05$. $\gamma = 0.13$				
(**) Significant association (χ^2), $p < 0.05$. $\gamma = -0.11$				

Table 7: Level of Agreement on Informational Usefulness based on the Evaluation of usage Experience.

		Bad or Very Bad	Average	Good or Very Good
With this type of tool, it is easier to be informed (*)	Disagree	47,5%	30,0%	153,6%
	Indifferent	25,4%	28,1%	21,5%
	Agreed	27,1%	41,8%	62,8%
	TOTAL	100,0%	100,0%	100,0%
This type of tool encourages the spread of hoaxes or fake news (**)	Disagree	11,9%	10,7%	22,9%
	Indifferent	22,0%	28,8%	28,3%
	Agreed	66,1%	60,5%	48,7%
	TOTAL	100,0%	100,0%	100,0%
Source: Own elaboration.				
(*) Significant association (χ^2), $p < 0.05$. $\gamma = 0.39$				
(**) Significant association (χ^2), $p < 0.05$. $\gamma = -0.26$				

The results related to prior experience with generative AI reveal interesting insights. To adequately characterise the students who reported having used ChatGPT and evaluated their experience with it, a multiple segmentation analysis was conducted using some of the variables previously identified as relevant. The technique employed (Chaid method) offers the advantage of graphically and intuitively classifying individuals into different groups based on the association of their personal characteristics with the evaluation of their experience (in this case, the dependent variable). The resulting segmentation tree (Chart 4) identifies distinct groups of students based on the most significant independent

variables that characterise them. In this case, these variables are the field of study of their degree programme, grouped age, and level of knowledge about how ChatGPT functions.

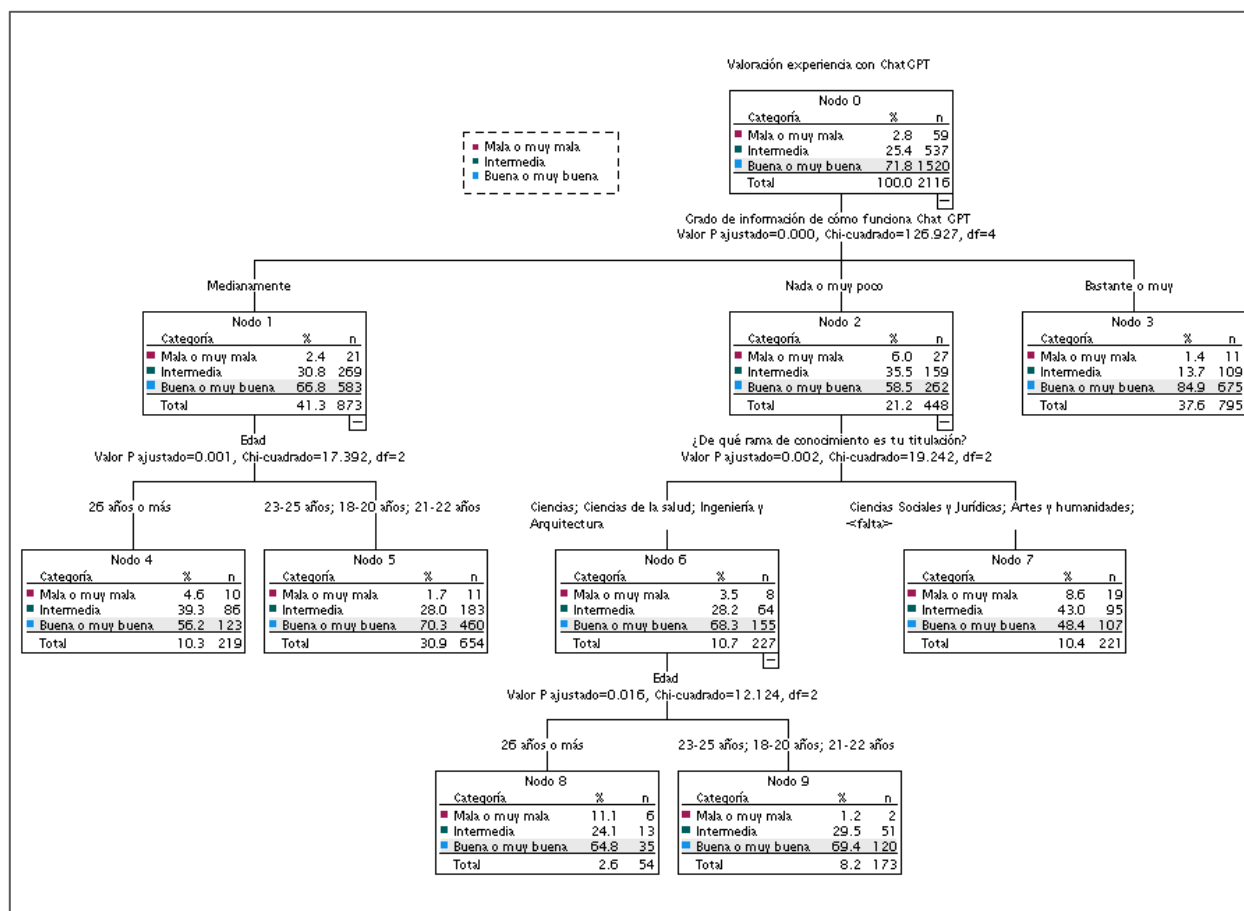


Chart 4: Multiple Segmentation Analysis Based on the Evaluation of ChatGPT Usage Experience.

Source: Own elaboration.

Risk estimate (incorrectly classified cases): 0.28. Growth method: Chaid.

- In this way, the resulting segmentation tree identifies six distinct groups of students (referred to as terminal nodes), characterised by the following features:
- Node 3 (37.6% of the sample): Fairly or very well-informed students.
- Node 4 (10.3% of the sample): Moderately informed students aged 26 or older.
- Node 5 (30.9% of the sample): Moderately informed students aged 18 to 25.
- Node 7 (10.4% of the sample): Students with very little or no knowledge, studying Social Sciences, Legal Sciences, or Arts and Humanities.
- Node 8 (2.6% of the sample): Students with very little or no knowledge, studying Sciences, Health Sciences, or Engineering and Architecture, aged 26 or older.
- Node 9 (8.2% of the sample): Students with very little or no knowledge, studying Sciences, Health Sciences, or Engineering and Architecture, aged 18 to 25.

The percentage of students who rate their overall experience with ChatGPT as good or very good is nearly 72%. However, when comparing these percentages across each of the terminal nodes, the lowest percentage is observed in Node 7 (48%), comprising students who report being very little or not at all informed. This is followed by Node 4 (56%), which includes students with an intermediate level of knowledge. In conclusion, to improve the user experience, it is essential to focus on increasing the level of knowledge, particularly among students in the fields of Social and Legal Sciences and Arts and Humanities, as well as those over 26 years old who have an intermediate level of knowledge. Students from other fields of study, despite having limited information, rate their experience as good or very good in 65% to 70% of cases, depending on their age group. As noted, these fields show better predisposition and more favourable evaluations of the various aspects related to the use of ChatGPT.

5. Thematic Analysis and Evaluation of Responses Generated by ChatGPT

As previously mentioned, the thematic analysis of ChatGPT's responses was conducted by comparing the information provided with official documentation available on the 11-M attacks, specifically the article *"La credibilidad de la fuente*

oficial. La secuencia 11-14 M” (Martínez-Solana, 2005) and the book “Los atentados del 11-M. Cuatro días que cambiaron la historia de España” (Martínez-Solana; Manchón-Campillo, 2022). In a second phase, the generative AI's responses were evaluated using the five semantic differential scales included in the student questionnaire: Incomplete/Complete, Doubtful/Reliable, Biased/Unbiased, Obsolete/Updated, and Incorrect/Correct. The following subsections, corresponding to each of the scales, present detailed examples of ChatGPT's responses to specific questions, along with justifications for the scores assigned. The complete set of scores assigned to the questions (distinguishing between those without a prompt and those with a prompt, marked with letters A and B, respectively) is summarised in Table 8. The full texts of the responses are included in Appendix I.

Incomplete – complete scale

The first question, posed without any additional considerations, is direct and straightforward: “What happened on 11 March 2004?” (1A). ChatGPT's response is:

"On that day, in 2004, a devastating terrorist attack took place in Madrid, the capital of Spain. Several commuter trains were targeted in a series of coordinated explosions, resulting in the deaths of 191 people and injuring more than 2,000. The attack, also known as the 11 March attacks or simply 11-M, was carried out by a group of Islamic extremists seeking to avenge Spain's involvement in the Iraq War, which had begun in March 2003. The attacks had a profound impact on Spanish society and the international community."

The response is rated with a score of 4 on a scale from 0 to 10 for the semantic pair incomplete/complete, as it omits important details such as the number of trains (four commuter trains in the Madrid suburban rail network), the specific locations (various points along the railway line connecting the Alcalá de Henares and Madrid-Atocha stations, known as "El Corredor del Henares"), the exact time (between 07:37 and 07:39), and the fact that there were a total of ten explosions. This basic information is documented in the *Boletín Oficial de las Cortes Generales* of July 14, 2005 (No. 242) (Martínez-Solana; Manchón-Campillo, 2022).

Additionally, the response is further evaluated as tending towards "incomplete" because it includes the following statement:

"11 March marked a turning point in the perception of the terrorist threat in Spain and had significant political repercussions. The government at the time, led by the Popular Party, initially attributed the attacks to the Basque separatist organisation ETA, which caused significant controversy. However, it was later discovered that the attack was carried out by a radical Islamist group linked to Al-Qaeda."

The text does not specify that the “political repercussions” mentioned by ChatGPT refer to the defeat of the Popular Party in the general elections held on 14 March 2004, after eight years in government.

Next, a new question (1B) was posed using a prompt that incorporated the criteria of role definition, source origin and type, context, and citations. The prompt read as follows:

"Act as a journalist and clearly explain, providing precise data, what happened in Madrid on 11 March 2004. Include information from various media outlets with varying ideological perspectives and place the events within the social and political context of the time. Cite the sources for your response."

The response receives a score of 8, placing it closer to the "complete" end of the incomplete/complete scale, as it includes the basic information about the question and cites three sources: BBC News, *El País*, and ABC. However, from a journalist's perspective, it lacks details such as the specific outcome of the 14 March elections. ChatGPT's response states:

"On 14 March 2004, in the general elections, the Spanish Socialist Workers' Party (PSOE), led by José Luis Rodríguez Zapatero, won the election with a clear majority. The 11-M attacks are considered to have had a significant impact on the electoral outcome, as many voters expressed dissatisfaction with the way the Popular Party government handled the situation."

To provide greater informational accuracy, the information should have included the detail that the PSOE won with 164 seats, while the PP obtained 148 seats.

Doubtful - reliable scale

Question 4A, supposedly posed by a person untrained in the use of generative AI, was as follows: “When were the perpetrators of the attack arrested?” ChatGPT's response primarily focuses on the timeline, with references to the days following the event:

"On 13 March 2004, just two days after the attacks, Spanish police arrested several suspects in connection with the attacks."

At the end of the paragraph, it also says:

"It should be noted that some of those involved in the 11-M attacks died in a confrontation with the police in April 2004,"

a few days after the attacks, when the authorities were trying to capture them in a suburb of Madrid."

The response receives a score of 4 on a scale of 0 to 10. To be more reliable, the information should specify that on 3 April 2004, the police located seven of the perpetrators of the attack in an apartment in Leganés (Madrid). It should also mention that they died after committing suicide by triggering an explosion, which also killed Sub-Inspector Francisco Javier Torronteras. These details are included in the Ruling of Investigative Proceedings 20/2004. (Martínez-Solana; Manchón-Campillo, 2022). On the other hand, in question 4B, ChatGPT is asked to provide an example of how to create a good prompt to obtain *"detailed, reliable, complete, and up-to-date information about the perpetrators of the attack that took place in Madrid on 11 March 2004."* ChatGPT provides several example questions, and we pose the third one to the chatbot: *"What were the connections of the perpetrators of the attack to international jihadist groups, such as Al-Qaeda?"*, corresponding to question 4B.

On this occasion, the response provided by the chat speaks of the perpetrators being linked to *"international jihadist groups, specifically the Al-Qaeda terrorist network"*.

This is completed in later paragraphs:

"Shortly after the attacks, a group calling itself <Brigadas Abu Hafs al-Masri> , allegedly linked to Al-Qaeda, claimed responsibility for the attacks in the name of Islamic jihad."

The response should include important information about the direct perpetrators, as after months of investigation, the Spanish government, with the collaboration of international organisations, determined that the perpetrators were linked to Al-Qaeda. However, Al-Qaeda itself was not the group directly responsible for carrying out the massacre. According to the Judgement of the National Court, the perpetrators of the attacks were ten individuals: Sarhane Ben Abdelmajid Fakheth, Jamal Ahmidan (known as "El Chino"), Mohamed Oulad Akcha, Rachid Oulad Akcha, Abdennabi Kounjaa, Asrih Rifaat Anouar, Allekema Lamari, Jamal Zougham, Otman el Gnaoui, and another unidentified individual. However, to this day, the mastermind behind the attack remains unknown (Martínez-Solana; Manchón-Campillo, 2022). The score obtained in this case is 6 on a scale of 0 to 10, as it continues to reproduce dubious information.

Biased - unbiased scale

Following the general question about 11-M posed at the beginning of the sequence, the inquiry continued with: *"Who perpetrated the attack?"* (Question 2A). The response received a score of 4 on a scale of 0 to 10, leaning towards a biased position. The first paragraph includes the following text:

"Initially, after the attacks on 11 March 2004 in Madrid, the Spanish government, led by the Popular Party, attributed the attacks to the Basque separatist organisation ETA (Euskadi Ta Askatasuna). This was due to the long-standing conflict between the Spanish government and ETA, which had led to numerous acts of violence in the past. However, as the investigation progressed, evidence emerged pointing to the involvement of a radical Islamist group linked to Al-Qaeda in the attacks. Traces of explosives and other indicators were found that did not match ETA's modus operandi. Additionally, connections to international terrorist cells were uncovered, and several individuals of Moroccan origin were identified as key suspects."

Based on the documentation consulted and according to anti-terrorism expert Mariano Rayón, on the day of the attack, the Arab newspaper *Al Quds Al Arabi*, published in London, reported that the terrorist group linked to the Al-Qaeda network, *Brigades of Abu Hafs Al-Masri*, had claimed responsibility for the attacks in the name of the international jihad. On 12 March, the Basque newspaper *Gara* and the Basque television network ETB shared a statement in which the terrorist group ETA denied involvement in the 11-M attacks: *"ETA is not responsible for yesterday's attacks."* This claim had already been made to the media on the morning of 11 March by the spokesperson for the Euskal Herritarrok coalition, Arnaldo Otegi.

Finally, on 13 March, an anonymous phone call informed authorities of the existence of a video left at the M-30 mosque in Madrid. In the video, a young man speaking Arabic with a Moroccan accent claimed responsibility for the attacks on behalf of the military wing of *Ansar Al-Qaeda in Europe*, Abu Dujan Al Afghani. This claim was further confirmed on 3 April via a fax sent to the newsroom of the newspaper ABC. This critical information regarding the hours immediately following the attacks, which shaped the determination of responsibility, should have been included in ChatGPT's response, given its relevance to what ChatGPT itself refers to as the *"political repercussions"* of the attacks (Martínez-Solana; Manchón-Campillo, 2022).

In question 2B, the prompt criterion used involves requesting the source of the information and specifying the format. The prompt is expressed as follows: *"Explain who the perpetrators of the 11-M attack were according to different media outlets. Provide the various explanations organised in a table so that the information can be compared."* In its response, after providing the table, ChatGPT adds the following:

"It is important to note that, while the 11-M attacks were perpetrated by a radical Islamist cell and not by ETA, there was initially some confusion and speculation in the media, particularly during the first hours and days following the"

attacks. As the investigation progressed, it was concluded that jihadist terrorism was responsible for the attacks."

By stating that *"initially there was some confusion and speculation in the media,"* ChatGPT fails to cite any specific media outlet that supported the theory attributing the attack to ETA. As a result, the response remains biased, though it receives a score of 8 on a scale of 0 to 10 based on the information provided. It is both significant and paradoxical that, at the end of the text, ChatGPT includes the following recommendation without offering any justification: *"Remember that objectivity and accuracy in journalism are fundamental, and it is essential to consult multiple sources to gain a complete and precise understanding of events."*

Obsolete - updated scale

In response to the direct question 5A, *"Point out the most important figures of 11-M,"* ChatGPT provides the following text, mentioning four individuals:

1. *"Jamal Zougam: One of the main suspects and accused individuals in the trial for the 11-M attacks. He was identified as one of the direct perpetrators and was convicted for his involvement in the attacks."*
2. *Serhane Ben Abdelmajid Fakhet: Also known as "El Tunecino" ("The Tunisian"), he was identified as one of the leaders of the radical Islamist group that carried out the attacks. He is believed to have been one of the key organisers of the attacks and died in the explosion at the Leganés apartment."*
3. *Mustafá Setmarián: Known as Abu Musab al-Suri, he was a Syrian jihadist theorist who lived in Spain for a time and was involved in jihadist activities. Although he was not directly implicated in the 11-M attacks, his presence in Spain and his influence on the jihadist environment are relevant for understanding the context."*
4. *José Luis Rodríguez Zapatero: Leader of the Spanish Socialist Workers' Party (PSOE) at the time and became the President of the Spanish Government after winning the general elections held three days after the attacks."*

The response does not mention key figures related to the 11-M attacks, such as José María Aznar, who was the president of the government at the time, or José Emilio Suárez Trashorras, an accomplice in the attack who provided the explosives to the terrorist group. The evaluation assigns a score of 4 on the obsolete/updated scale.

Question 5B, *"Create a table with the most important figures of 11-M, justifying their role with a journalistic approach and aimed at someone who wants to be well-informed on the subject,"* mentions the key perpetrators of the attacks (Jamal Zougam, "El Chino," Serhane Ben Abdelmajid, "El Tunecino," and Rabel Osman Sayed Ahmed, known as "Mohamed El Egipcio"), the jihadist terrorism expert Fernando Reinares, journalist Iñaki Gabilondo, the prime ministers José María Aznar and José Luis Rodríguez Zapatero, and the Abu Hafs al-Masri Brigades.

The response to this question received a score of 7 on a scale from 0 to 10. However, it is important to note that at the end, ChatGPT states: *"My responses are generated based on the knowledge and data acquired up to my cut-off date in September 2021. I do not have internet access to verify updated information. While I strive to provide accurate and reliable answers, some specific details may change or evolve over time."* This comment is not relevant in this case, as the available information, such as the judicial ruling on the 11-M trial dated 31 October 2007, predates 2021.

Incorrect - correct scale

In question 7A, *"How is Aznar's handling of the 11-M attack evaluated?"*, we asked ChatGPT to summarise opinions regarding the decisions and positions taken by the government following the attack. The AI begins by summarising that:

"It is a topic that generates diverse opinions and may vary depending on different political and social perspectives."

It then presents the critical arguments regarding Aznar's handling of the attack:

"Some critics have questioned how the Aznar government managed information and communication in the days following the attacks. Initially, an attempt was made to attribute the attacks to ETA, despite evidence pointing to an Islamist origin. This confusion and the subsequent correction of the information led to criticism about transparency and the handling of the crisis."

Information on Spain's participation in the Iraq War is also included, a decision that had generated widespread debate and disagreement in part of the Spanish population:

"Moreover, some have criticised Aznar's foreign policy, including his alignment with the United States-led coalition in the Iraq War. The 11-M attacks were considered by some as a response to Spain's support for the invasion of Iraq, which sparked a debate about the extent to which the government's actions may have increased the terrorist threat in Spain."

In the penultimate paragraph, ChatGPT presents the perspective of "defenders of Aznar":

"On the other hand, some defenders of Aznar have highlighted the efforts made to combat jihadist terrorism during his government, including the arrest and dismantling of terrorist cells in Spain. It has also been argued that attributing

political responsibility for the attacks to Aznar's government oversimplifies the complexity of the factors that contributed to the attacks."

In the final paragraph of the answer, ChatGPT concludes that:

"In summary, the evaluation of Aznar's handling of 11-M is a controversial and polarising topic, with divergent opinions on transparency in communication, foreign policy, and the actions taken to prevent and respond to the terrorist threat in Spain."

In the evaluation of the controversy, the response does not include verifiable information to support the argument presented by Aznar's defenders, resulting in a score of 4 on the incorrect/correct scale (0 to 10).

Finally, in question 7B, the prompt requests a detailed explanation of how Aznar managed the political and social aspects of the 11-M attacks, for someone seeking to be well-informed on the subject with a neutral perspective and specifies that the sources of information should be included.

For the first time, the response is considered correct, as ChatGPT successfully cites sources of information, predominantly referencing Spanish media outlets that were key sources of information on the topic. In this case, the response is rated an 8.

"Initial Attribution to ETA: After the attacks, Aznar's government quickly attributed responsibility to ETA, the Basque separatist terrorist organisation. This decision was later criticised, as it was revealed that there was no evidence to support this claim and that, in fact, the responsibility lay with a jihadist cell. (El País, ABC.es)"

"Communication Crisis: The initial attribution to ETA sparked controversy and distrust within Spanish society, as the majority of the public doubted the accuracy of the information provided by the government. This communication crisis intensified when investigations pointed towards a possible jihadist connection. (El Mundo, La Vanguardia)"

Table 8: Expert Evaluation of the Answers Provided by ChatGPT.

Questions Asked		Ratings (0 to 10)				
		Incomplete Complete	Doubtful Reliable	Biased Unbiased	Obsolete Updated	Incorrect Correct
1A	What happened on 11-M in 2004?	4	4	6	4	4
1B	Act as a journalist and explain to me clearly and with precise data what happened in Madrid on 11 March 2004. Include information from various media sources with diverse ideological orientations and place the events in the social and political context of the time. Provide the sources of your answer.	8	6	4	9	6
2A	Who were responsible for the attack?	4	4	4	5	5
2B	Explain who was responsible for the 11-M attack according to the several media outlets. Give the various explanations in a table so I can compare the information.	5	7	8	6	7
3A	What role did ETA play in 11-M?	4	4	4	5	4
3B	I want you to provide a detailed contextualisation with arguments about the role ETA played in the 11-M attacks. I would like you to include the sources from which you obtain the information in all cases.	6	7	7	8	8
4A	When were the perpetrators of the attack arrested?	7	4	6	7	3
4B	"Give me examples of some <i>prompts</i> that I can use to obtain detailed, reliable, complete and up-to-date information about the perpetrators of the attack that took place in Madrid on 11 March 2004." What were the connections of the authors of the attack with international jihadist groups, such as Al-Qaeda?	8	6	7	8	6
5A	Point out the important characters of 11-M	4	4	4	4	4
5B	Create a table with the most important figures of 11-M, explaining their roles from a journalistic approach and for someone who wants to be well-informed on the subject.	7	6	5	7	5
6A	How was President Aznar have related to 11-M?	4	4	4	3	3
6B	Provide me with an evaluation in the journalistic format of an opinion article, aimed at a university student, on how Aznar's management of the 11-M attacks was handled.	7	7	8	7	8
7A	How is Aznar's management of 11-M evaluated?	4	4	3	4	4
7B	Specify how the political and social management was handled by José María Aznar, who was the President of the Spanish Government at the time of the 11-M attacks in 2004. Prepare the response for someone who wants to be well-informed on the topic, with a neutral perspective. Include the sources from which you extract the information; use more than five different sources.	7	7	8	7	8

Table 9 shows the summary of expert scores. It can be seen that ChatGPT performs noticeably better when questions are posed using well-constructed prompts, with a net improvement of between 2 and 3 points across the five scales considered. However, it is important to note that even when using prompts, maximum scores are not achieved, as there remains room for improvement in the responses. This could potentially be addressed with the latest versions of generative AI (4.0) and more extensive training procedures on the specific topic. When comparing students' subjective perceptions of ChatGPT, as reflected in the online questionnaire, with the results of the thematic analysis of questions without prompts (assumed to be posed by someone untrained in the use of generative AI), it is evident that students rate ChatGPT's responses more positively across all five scales considered. However, when compared to expert evaluations of questions posed with prompts, it becomes clear that students' ratings are slightly lower than expert evaluations across all scales, except for the Obsolete-Updated scale, where, according to expert judgement, ChatGPT shows a significant improvement.

Table 9: Expert's Average Ratings of the Responses Provided by ChatGPT (scales from 0 to 10).

Average Ratings											
	Incomplete	Complete	Doubtful	Reliable	Biased		Unbiased	Obsolete	Updated	Incorrect	Correct
All the questions	5.64		5.29		5.57		6.00		5.36		
Questions without <i>prompt</i>	4.43		4.00		4.43		4.57		3.86		
Questions with <i>prompt</i>	6.83		6.50		6.50		7.50		6.67		
improvement with <i>prompt</i>	2.40		2.50		2.07		2.93		2.81		

6. Discussion and Conclusions

The analysis of the online survey conducted with university students reveals that greater knowledge about generative AI, and especially greater exposure to it, increases the likelihood of having a more positive perception of ChatGPT. These results confirm the hypotheses posed in the research and demonstrate that students' perceptions of ChatGPT are influenced by their level of knowledge about the tool and their prior experience using it. Additionally, the findings show that students who perceive themselves as very well informed about how ChatGPT works and those who report having significant experience with it (as opposed to those with moderate or very little experience) are more likely to agree that such tools facilitate access to information. Conversely, students who report being not at all or very little informed about ChatGPT's functionality, and who have limited experience using it, are more inclined to agree that these tools promote the spread of fake news or misinformation.

This self-perception is also associated with sociodemographic and academic variables. Men score higher than women on both indicators, as do students in the field of Engineering and Architecture. Regarding age, the level of knowledge is higher among those aged 26 and over, while experience with ChatGPT use is most prominent among younger students aged 18 to 22. The analysis of the evaluation of the information provided by ChatGPT also reveals statistically significant means' differences (on a scale from 0 to 10) across the semantic differential scales used. Students who report a poor or very poor experience with ChatGPT rate the information as tending towards incomplete (2.91), doubtful (2.40), biased (3.33), obsolete (4.45), and incorrect (2.42). Conversely, those who report a good or very good experience rate the information as tending towards complete (7.17), reliable (6.26), unbiased (6.04), updated (6.75), and correct (6.53).

The average evaluations of the information provided by ChatGPT, based on the sociodemographic variables of gender and social class, show that the indicators receiving the lowest ratings (around 5) are **biased/unbiased** and **doubtful/reliable**. A similar trend is observed in relation to academic variables, such as the type of study programme and the field of study. These findings align with other studies that have focused on analysing university students' perceptions. In research conducted by **Grassini** (2023) on the use of ChatGPT in education, several challenges were identified, including issues with accuracy, reliability, and biases in the information provided. Similarly, **Ngo** (2023) highlights that students acknowledge ChatGPT's ability to deliver information across various domains but recognise certain issues, such as its inability to evaluate the quality of information and its failure to accurately cite sources.

The multiple segmentation analysis identifies distinct profiles of university students based on their prior experience with ChatGPT. This is particularly valuable, as university students cannot be regarded as a homogenous group. Based on their level of knowledge, age, and field of study, the six groups identified, provide a basis for targeting the actions and policies that universities must undoubtedly implement to expand learning and usage of generative AI. The findings suggest a need to focus efforts on older students (26 years or older) who have some knowledge of ChatGPT, as well as on those who report being very little or not at all informed and are enrolled in the fields of Social and Legal Sciences or Arts and Humanities. By addressing these specific groups, universities can tailor their strategies to improve awareness, adoption, and effective use of generative AI.

The article also delves into the evaluation of the information provided by ChatGPT, acknowledging that this is a crucial factor in determining whether the technology will be adopted in educational and professional development practices. Using the same semantic differential scales by Osgood that were incorporated into the student questionnaire and taking the topic of the 11-M attacks in Madrid as an example, it is observed that expert evaluations by individuals knowledgeable about the subject rate responses to direct questions without prompts as more likely to be incomplete, doubtful, biased, obsolete, and incorrect. However, this trend is reversed when evaluating responses to questions posed with prompts. This leads to the conclusion that students must be trained in the effective use of generative AI, focusing on dimensions such as those analysed in this article: incomplete/complete, doubtful/reliable, biased/unbiased, obsolete/updated, and incorrect/correct.

Additionally, it is worth noting a tendency to achieve better results as question formulation improves. This corresponds with existing literature emphasising the importance of AI training, which highlights that asking the same question without context versus with additional context yields different responses, due to the probabilistic model on which ChatGPT is based (**García-Peñalvo**, 2023). When the AI is prompted to adopt the role of a journalist, summarising and organising information into tables, or explicitly asked to cite the sources used, the evaluations of its responses also improve. However, on some occasions, ChatGPT provides assessments without including sources that allow its argumentation to be verified.

These conclusions align with findings by **Gutiérrez-Caneda et al. (2023)** regarding the analysis of content provided by ChatGPT. The authors highlighted the concern of participants in their study about the ease with which ChatGPT can include false information, particularly because it presents it in a format and style that makes it appear relevant, credible, and reliable. **Grassini (2023)** also notes that, in some cases, the information may be politically biased or include biases related to race, gender, or equity. Additionally, **Pavlik (2023)**, in an essay based on questions posed to ChatGPT and the responses obtained, highlights a notable example where ChatGPT was asked, "What should journalism professors teach about AI?" The response provided by ChatGPT itself acknowledged both the advantages of AI—such as its ability to quickly process large volumes of data—and its potential drawbacks, including the risk of incorporating biases and errors into reports.

The main problems identified are related to the generation of **incorrect information** and **biased information**. Incorrect information is sometimes associated with the phenomenon known as the "**hallucination effect**", where AI technologies fabricate information (**Vargas-Murillo et al., 2023**). Additionally, generative models still face significant limitations in their capabilities, such as a lack of conceptualisation and generalisation of knowledge, as well as their inefficiency in terms of energy, time, and computational resources required to retrain the models. These issues, among others, have heightened concerns over ethical and political matters, such as **transparency, privacy, and accountability**, making it necessary to carefully evaluate how these models are used and how they can be regulated to ensure ethical and responsible utilisation (**Lipton, 2018**). Experts advocate for continuous oversight, emphasising the application of ethical codes at different levels of information production and among the various professionals using these tools (**Noain-Sánchez, 2022**). The European Union has also taken its first steps towards regulating the use of Artificial Intelligence technologies, with the approval of the **AI Regulation (Unión Europea, 2024)**.

The need for users to be aware of these risks and take measures to mitigate them has been emphasised in various fields. The first step involves the AI industry itself, which must aim to address and correct some of these issues through updates. A significant advancement was made with the launch in May 2024 of **GPT-4o**, which is connected to the internet, provides updated information and is no longer limited to 2021 as was the case with previous versions. Similarly, Google has released **Gemini**, a replacement for the previous Bard, offering greater capacity to interpret and generate text, images, and audio. The growing concern over the risks associated with generative AI is reflected in initiatives such as the recent **repository of identified risks** compiled by MIT (**Slattery et al., 2024**) or the **PRISM model** proposed by the World Economic Forum to promote responsible AI (**World Economic Forum, 2024**).

The second, and most important, measure involves analysing the possibilities and limitations of ChatGPT as an educational tool and addressing issues such as the **digital divide**, the development of **critical skills learning** (**Baidoo-Anu; Ansah, 2023; Ribera; Díaz, 2024**), and **academic integrity**. One of the main challenges lies in the speed at which these innovations must be assessed and incorporated (**García-Peñalvo et al., 2024**), as well as the need to develop an **ethical framework** and guidelines to guide their implementation and use in Higher Education (**Abas et al., 2023**). To achieve this, more **interdisciplinary research** should be promoted, involving participants from various age groups and professional contexts (**Chu et al., 2022**), as well as experiences that provide practical insights into the use of ChatGPT in Higher Education, as explored in this article. The findings from such studies will facilitate the development of training programmes that foster a stronger and more effective relationship and integration between AI, students, educators, and information professionals.

References

- Abas, Mohd Azman; Arumugam, Sathi Eswari; Yunus, Melor Md; Rafiq, Karmila Rafiqah M. (2023).** "ChatGPT and Personalized Learning: Opportunities and Challenges in Higher Education". *International Journal of Academic Research in Business and Social Sciences*, v. 13, n. 12, pp. 3536-3545. <https://doi.org/10.6007/IJARBS/v13-i12/20240>
- Arroyo Menéndez, Millán; Finkel Morgenstern, Lucila. (2019).** "Encuestas por Internet y nuevos procedimientos muestrales". *Panorama Social*, v. 30, pp. 41-53. <https://hdl.handle.net/20.500.14352/13911>
- Baidoo-Anu, David; Ansah, Leticia Owusu. (2023).** "Education in the Era of Generative Artificial Intelligence (AI): Understanding the Potential Benefits of ChatGPT in Promoting Teaching and Learning". *Journal of AI*, v. 7, n. 1, pp. 52-62. <https://doi.org/10.61969/jai.1337500>
- Bennett, W. Lance; Lawrence, Regina G.; Livingston, Steven. (2006).** "None Dare Call It Torture: Indexing and the Limits of Press Independence in the Abu Ghraib Scandal". *Journal of Communication*, v. 56, n. 3, pp. 467-485. <https://doi.org/10.1111/j.1460-2466.2006.00296.x>
- Bringula, Rex. (2023).** "What Do Academics Have to Say About ChatGPT? A Text Mining Analytics on the Discussions Regarding ChatGPT on Research Writing". *AI and Ethics*. <https://doi.org/10.1007/s43681-023-00354-w>
- Centeno-Martí-n, Héctor; Toledano-Buendí-a, Samuel; Ardèvol-Abreu, Alberto. (2023).** "Who interacts with communicative AI and what attitudes are displayed toward it? Sociodemographic, personality, and futurephobia variables". *Profesional de la información*, v. 32, n. 5, pp. e320502. <https://doi.org/10.3145/epi.2023.sep.02>

- Chu, Shih-Ting; Hwang, Gwo-Jen; Tu, Yun-Fang.** (2022). "Artificial Intelligence-Based Robots in Education: A Systematic Review of Selected SSCI Publications". *Computers and Education: Artificial Intelligence*, v. 3, pp. 100091. <https://doi.org/10.1016/j.caeai.2022.100091>
- Das, Soumya Ranjan; Madhusudan, J V.** (2024). "Perceptions of Higher Education Students towards ChatGPT Usage". *International Journal of Technology in Education*, v. 7, n. 1, pp. 86-106. <https://doi.org/10.46328/ijte.583>
- Díaz-de-Rada, Vidal.** (2022). "Strategies to Improve Response Rates to Online Survey". *Papers*, v. 107, n. 4, pp. e3073. <https://doi.org/10.5565/rev/papers.3073>
- Flores-Viva, Jesús-Miguel; García-Peñalvo, Francisco-José.** (2023). "Reflections on the Ethics, Potential, and Challenges of Artificial Intelligence in the Framework of Quality Education (SDG4)". *Comunicar*, v. 31, n. 74, pp. 35-44. <https://doi.org/10.3916/C74-2023-03>
- García-Peñalvo, Francisco-Javier; Llorens-Largo, Faraón; Vidal, Javier.** (2024). "La nueva realidad de la educación ante los avances de la inteligencia artificial generativa". *RIED-Revista Iberoamericana de Educación a Distancia*, v. 27, n. 1, pp. 9-39. <https://doi.org/10.5944/ried.27.1.37716>
- García-Peñalvo, Francisco-José.** (2023). "The perception of Artificial Intelligence in educational contexts after the launch of ChatGPT: Disruption or panic?". *Education in the Knowledge Society (EKS)*, v. 24, pp. e31279. <https://doi.org/10.14201/eks.31279>
- García-Sánchez, Omar-Vicente.** (2023). "Uso y percepción de ChatGPT en la educación superior". *Revista de Investigación en Tecnologías de la Información*, v. 11, n. 23, pp. 98-107. <https://doi.org/10.36825/RITI.11.23.009>
- García, Noelia.** (2023). *¿Es ChatGPT el fin del pensamiento crítico en las aulas?* El Economista. 22 de febrero. <https://www.economista.es/tecnologia/noticias/12159203/02/23/Es-Chat-GPT-el-fin-del-pensamiento-critico-en-las-aulas.html>
- Giray, Louie.** (2023). "Prompt Engineering with ChatGPT: A Guide for Academic Writers". *Annals of Biomedical Engineering*, v. 51, n. 12, pp. 2629-2633. <https://doi.org/10.1007/s10439-023-03272-4>
- Grassini, Simone.** (2023). "Shaping the Future of Education: Exploring the Potential and Consequences of AI and ChatGPT in Educational Settings". *Education Sciences*, v. 13, n. 7, pp. 692. <https://doi.org/10.3390/educsci13070692>
- Grizzle, Alton; Wilson, Carolyn; Tuazon, Ramon; Cheung, Chi Kim; Lau, Jesús; Fischer, Rachel; Gordon, Dorothy; Akyempong, Kwame; Singh, Jagtar; Carr, Paul R.; Stewart, Kristine; Tayie, Samy; Suraj, Olunifesi; Jaakkola, Maarit; Thésée, Gina; Gulston, Curmira; Andzongo-Menyeng, Blaise-Pascal; Zibi-Fama, Paul-Alain.** (2023). *Ciudadanía alfabetizada en medios e información: pensar críticamente, hacer clic sabiamente*. UNESCO. <https://unesdoc.unesco.org/ark:/48223/pf0000385119>
- Gutiérrez-Caneda, Beatriz; Vázquez-Herrero, Jorge; López-García, Xosé.** (2023). "AI application in journalism: ChatGPT and the uses and risks of an emergent technology". *Profesional de la información*, v. 32, n. 5, pp. e320514. <https://doi.org/10.3145/epi.2023.sep.14>
- Lample, Guillaume; Conneau, Alexis.** (2020). "Cross-lingual Language Model Pretraining". *arXiv preprint arXiv:1901.07291*. <https://doi.org/10.48550/arXiv.1901.07291>
- Li, Jiwei; Monroe, Will; Jurafsky, Dan.** (2021). "Understanding Neural Networks through Representation Erasure". *arXiv:1612.08220*. <https://doi.org/10.48550/arXiv.1612.08220>
- Lipton, Zachary C.** (2018). "The mythos of model interpretability". *Queue*, v. 16, n. 3, pp. 31-57. <https://doi.org/10.1145/3236386.3241340>
- Márquez, Javier.** (2023). "GPT4: así es la IA más avanzada de OpenAI, cómo funciona y todas las novedades." *Xataka*, 17 mayo, 2023. <https://www.xataka.com/nuevo/gpt-4-que-cuando-sale-como-funciona-toda-informacion>
- Martínez-Solana, Yolanda.** (2005). "La credibilidad de la fuente oficial. La secuencia 11-14 M". *Estudios sobre el Mensaje Periodístico*, v. 11, pp. 65-92. <https://revistas.ucm.es/index.php/ESMP/article/view/ESMP0505110065A>
- Martínez-Solana, Yolanda; Manchón-Campillo, Felipe.** (2022). *Los atentados del 11-M: cuatro días que cambiaron la historia de España*. Ediciones Complutense. <https://www.ucm.es/ediciones-complutense/los-atentados-del-11m-cuatro-dias-que-cambiaron-la-historia-de-espana>
- Micó, Josep-Lluís; Casero-Ripollés, Andreu; García-Orosa, Berta.** (2022). "Platforms in Journalism 4.0: The Impact of the Fourth Industrial Revolution on the News Industry." In: *Total Journalism: Models, Techniques and Challenges*. Vázquez-Herrero, Jorge; Silva-Rodríguez, Alba; Negreira-Rey, María-Cruz; Toural-Bran, Carlos; López-García, Xosé (Eds.), pp. 241-253. Springer International Publishing. https://doi.org/10.1007/978-3-030-88028-6_18
- Mollick, Ethan R.; Mollick, Lilach.** (2023). "Assigning AI: Seven approaches for students, with prompts". *The Wharton School Research Paper*. <https://doi.org/10.2139/ssrn.4475995>

- Morath, Benedict; Chiriac, Ute; Jaszowski, Elena; Deiß, Carolin; Nürnberg, Hannah; Hörth, Katrin; Hoppe-Tichy, Torsten; Green, Kim.** (2024). "Performance and Risks of ChatGPT used in Drug Information: An Exploratory Real-World Analysis". *European Journal of Hospital Pharmacy*, v. 31, n. 6, pp. 491-497. <https://doi.org/10.1136/ejhp-2023-003750>
- Ngo, Thi Thuy An.** (2023). "The perception by university students of the use of ChatGPT in education". *International Journal of Emerging Technologies in Learning (IJET)*, v. 18, n. 17, pp. 4-19. <https://doi.org/10.3991/ijet.v18i17.39019>
- Noain-Sánchez, Amaya.** (2022). "Addressing the Impact of Artificial Intelligence on Journalism: The perception of experts, journalists and academics". *Communication and Society*, v. 35, n. 3, pp. 105-121. <https://doi.org/10.15581/003.35.3.105-121>
- Obenza, Brandon-Nacua; Salvahan, Alexa; Rios, Alexandra-Nicole; Solo, Althea; Alburo, Rea-Ahlee; Gabila, Rey-Jose.** (2024). "University Students' Perception and Use of ChatGPT: Generative Artificial Intelligence (AI) in Higher Education". *International Journal of Human Computing Studies*, v. 5, n. 12, pp. 5-18. <https://doi.org/10.5281/zenodo.10360697>
- Ocaña-Fernández, Yolvi; Valenzuela-Fernández, Luis-Alex; Garro-Aburto, Luzmila.** (2019). "Inteligencia artificial y sus implicaciones en la educación superior". *Propósitos y Representaciones*, v. 7, n. 2, pp. 536-568. <https://doi.org/10.20511/pyr2019.v7n2.274>
- Osgood, Charles E.; Suci, George J.; Tannenbaum, Percy H.** (1957). *La medida del significado*. Editorial Gredos. [https://www.uv.es/seoane/publicaciones/Osgood,%20Suci%20y%20Tannenbaum%201957%20La%20Medida%20del%20Significado%20\(1976\).pdf](https://www.uv.es/seoane/publicaciones/Osgood,%20Suci%20y%20Tannenbaum%201957%20La%20Medida%20del%20Significado%20(1976).pdf)
- Pavlik, John V.** (2023). "Collaborating With ChatGPT: Considering the Implications of Generative Artificial Intelligence for Journalism and Media Education". *Journalism & Mass Communication Educator*, v. 78, n. 1, pp. 84-93. <https://doi.org/10.1177/10776958221149577>
- Polo, Juan-Diego.** (2023). "ChatGPT: ¿Cómo obtener fuentes fiables?" *Whatsnew*, 14 marzo, 2023. <https://www.hatsnew.com/2023/03/14/chatgpt-como-obtener-fuentes-y-citas-fiables>
- Radford, Alec; Wu, Jeffrey; Child, Rewon; Luan, David; Amodei, Dario; Sutskever, Ilya.** (2019). "Language models are unsupervised multitask learners". *OpenAI Blog*, v. 1, n. 8, pp. 9. <https://d4mucfjksyvv.cloudfront.net/better-language-models/language-models.pdf>
- Ribera, Mireia; Díaz, Oliver.** (2024). *ChatGPT y educación universitaria: posibilidades y límites de ChatGPT como herramienta docente*. Universitat de Barcelona. IDP/ICE & Ediciones Octaedro. <https://diposit.ub.edu/dspace/bitstream/2445/206141/1/9788410054011.pdf>
- Rodríguez-Rivera, Paula; Camacho-Sánchez, Rubén; Rodríguez-Ferrer, José-Miguel; Manzano-León, Ana.** (2023). "Estudio sobre las percepciones y la aplicación educativa de ChatGPT en la Educación Superior de Galicia." In: *Tecnofilosofía: reflexión filosófica, inteligencia artificial y ciencia*. Bermúdez-Vázquez, Manuel; Sánchez-Cotta, Agustín (Eds.), pp. 211-223. Dykinson. <https://www.torrossa.com/en/resources/an/5713928#>
- Rodríguez, Eva.** (2023). "Estos son los siete errores más comunes al usar ChatGPT y otras inteligencias artificiales. Corregirlos es muy sencillo." *Genbeta*, 6 mayo, 2023. <https://www.genbeta.com/paso-a-paso/siete-errores-habituales-chatgpt-otras-herramientas-inteligencia-artificial>
- Russell, Stuart; Norvig, Peter.** (2021). *Artificial Intelligence: A Modern Approach*. Pearson. https://people.engr.tamu.edu/guni/csce421/files/AI_Russell_Norvig.pdf
- Sabzalieva, Emma; Valentini, Arianna.** (2023). *ChatGPT e inteligencia artificial en la educación superior: guía de inicio rápido*. UNESCO. https://unesdoc.unesco.org/ark:/48223/pf0000385146_spa
- Sánchez-Carrión, Juan-Javier; Segovia, José-Manuel; Sánchez-Meseguer, Paula.** (2012). "Las encuestas en internet." In: *Metodología de la investigación social: técnicas innovadoras y sus aplicaciones*. Arroyo, M.; Sádaba, I. (Eds.), pp. 79-108. Editorial Síntesis. <https://dialnet.unirioja.es/servlet/articulo?codigo=4571242>
- Sánchez-García, Pilar; Merayo-Álvarez, Noemí; Calvo-Barbero, Carla; Díez-Gracia, Alba.** (2023). "Desarrollo tecnológico español de la inteligencia artificial aplicada al periodismo: empresas y herramientas de documentación, producción y distribución de información". *Profesional de la información*, v. 32, n. 2, pp. e320208. <https://doi.org/10.3145/epi.2023.mar.08>
- Sart, Gamze.** (2023). "The Perception of the University Students of the Use of ChatGPT During the University Education: Case of Turkey." In: *Edulearn23 Proceedings*. pp. 8545-8545. IATED. <https://doi.org/10.21125/edulearn.2023.2269>
- Saz-Pérez, Fabio; Pizà-Mir, Bartolomé.** (2024). "Desafiando el estado del arte en el uso de ChatGPT en educación en el año 2023". *REIRE. Revista d'Innovació i Recerca en Educació*, v. 17, n. 1, pp. 1-13. <https://doi.org/10.1344/reire.44018>
- Segovia-García, Nuria.** (2023). "Percepción y uso de los chatbots entre estudiantes de posgrado online: Un estudio exploratorio". *Revista de Investigación en Educación*, v. 21, n. 3, pp. 335-349. <https://doi.org/10.35869/reined.v21i3.4974>

- Sier, Jessica.** (2022). "ChatGPT toma Internet por asalto, mala poesía y todo." *Financial Review*, 8 diciembre, 2022. <https://www.afr.com/technology/chatgpt-takes-the-internet-by-storm-bad-poetry-and-all-20221207-p5c4hv>
- Slattery, Peter; Saeri, Alexander K.; Grundy, Emily A. C.; Graham, Jess; Noetel, Michael; Uuk, Risto; Dao, James; Pour, Soroush; Casper, Stephen; Thompson, Neil.** (2024). *The AI Risk Repository: A Comprehensive Meta-Review, Database, and Taxonomy of Risks From Artificial Intelligence*. MIT AI Risk Repository. <https://doi.org/10.13140/RG.2.2.28850.00968>
- Stokel-Walker, Chris; Van Noorden, Richard.** (2023). "What ChatGPT and Generative AI Mean for Science". *Nature*, v. 614, n. 7947, pp. 214-216. <https://doi.org/10.1038/d41586-023-00340-6>
- Torres-Salinas, Daniel; Arroyo-Machado, Wenseslao.** (2023). *Manual de ChatGPT: aplicaciones en investigación y educación universitaria*. InfluScience Ediciones, Granada. <https://digibug.ugr.es/handle/10481/85145>
- Túñez-López, José-Miguel; Feiras-Ceide, César; Vaz-Álvarez, Martín.** (2021). "Impact of Artificial Intelligence on Journalism: Transformations in the Company, Products, Contents and Professional Profile". *Communication & Society*, v. 34, n. 1, pp. 177-193. <https://doi.org/10.15581/003.34.1.177-193>
- UC3M.** (2023). *Recomendaciones para la docencia con inteligencias artificiales generativas*. Universidad Carlos III de Madrid. <https://hdl.handle.net/10016/37989>
- Unión Europea.** (2024). *Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence and amending Regulations*. Official Journal of the European Union. July 12, 2024. <http://data.europa.eu/eli/reg/2024/1689/oj>
- Valova, Irena; Mladenova, Tsvetelina; Kanev, Gabriel.** (2024). "Students' Perception of ChatGPT Usage in Education". *International Journal of Advanced Computer Science & Applications*, v. 15, n. 1, pp. 466-473. <https://doi.org/10.14569/IJACSA.2024.0150143>
- Vargas-Murillo, Alfonso-Renato; Pari-Bedoya, Ilda; Guevara-Soto, Francisco-de-Jesús.** (2023). "Challenges and Opportunities of AI-Assisted Learning: A Systematic Literature Review on the Impact of ChatGPT Usage in Higher Education". *International Journal of Learning, Teaching and Educational Research*, v. 22, n. 7, pp. 122-135. <https://doi.org/10.26803/ijlter.22.7.7>
- Vaswani, Ashish; Shazeer, Noam; Parmar, Niki; Uszkoreit, Jakob; Jones, Llion; Gomez, Aidan N; Kaiser, Lukasz; Polosukhin, Illia.** (2017). "Attention Is All You Need." In: *Proceedings of the 31st International Conference on Neural Information Processing Systems*. Luxburg, Ulrike von; Guyon, Isabelle; Bengio, Samy; Wallach, Hanna; Fergus, Rob (Eds.), pp. 6000-6010. Curran Associates Inc. <https://dl.acm.org/doi/10.5555/3295222.3295349>
- Walters, William H.; Wilder, Esther Isabelle.** (2023). "Fabrication and errors in the bibliographic citations generated by ChatGPT". *Scientific Reports*, v. 13, n. 1, pp. 14045. <https://doi.org/10.1038/s41598-023-41032-5>
- White, Jules; Fu, Quchen; Hays, Sam; Sandborn, Michael; Olea, Carlos; Gilbert, Henry; Elnashar, Ashraf; Spencer-Smith, Jesse; Schmidt, Douglas C.** (2023). "A Prompt Pattern Catalog to Enhance Prompt Engineering with ChatGPT". *arXiv preprint arXiv:2302.11382*. <https://doi.org/10.48550/arXiv.2302.11382>
- World Economic Forum.** (2024). *AI for Impact: The PRISM Framework for Responsible AI in Social Innovation* [White Paper]. June 25. <https://www.weforum.org/publications/ai-for-impact-the-prism-framework-for-responsible-ai-in-social-innovation>
- Yusuf, Abdullahi; Pervin, Nasrin; Román-González, Marcos; Noor, Norah Md.** (2024). "Generative AI in education and research: A systematic mapping review". *Review of Education*, v. 12, n. 2, pp. e3489. <https://doi.org/10.1002/rev3.3489>