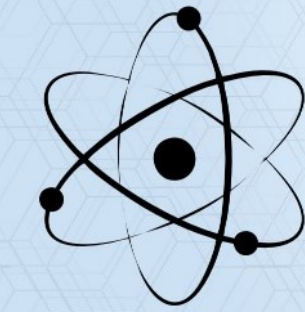




TheSci

International Science Journal[©]

SCIENCE-ENGINEERING-TECHNOLOGY

Scientific Journals

Unveiling the Future, Understanding the Present.

In every issue, dive with us into the latest scientific and technological innovations shaping our tomorrow. We explore the boldest advancements, analyze their implications today, and offer you a clear vision of how science and technology are transforming the world.



Issue 1, Volume 1, No. 01, January-June 2025



Review TheSci

ISSN: In process

Issue 1, Volume 1, No. 01, January-June 2025

TheSci

ISSN: In process



Issue 1, Volume 1, No. 01, January-June 2025

Do AIs Really Surpass Us? A Comparative Analysis of ChatGPT, Qwen, and DeepSeek in Key Human Skills.

*PhD. Alfredo Gutiérrez García

Universidad Tecnológica del Norte de Guanajuato

alfredo.gutierrez@utng.edu.mx
<https://orcid.org/0000-0002-8197-8144>

Received: 04-10-2025 | Accepted: 04-26-2025 | Published: 04-05-2025.

* Corresponding Author: alfredo.gutierrez@utng.edu.mx

How to cite this article: Gutiérrez Alfredo (2025). Do AIs Really Surpass Us? A Comparative Analysis of ChatGPT, Qwen, and DeepSeek in Key Human Skills. TheSci Review, 1 (1) 37-62. Quality Consulting Instituto de Educación Capacitación y Certificación de México. <https://ieccmexico.com/thesci>

Copyright (c). 2025 PhD. Alfredo Gutiérrez García; This is an open access article distributed under the terms of the Attribution 4.0 International ([CC BY-NC-SA 4.0](https://creativecommons.org/licenses/by-nc-sa/4.0/)) TheSci. Mexico Review / Vol. 1, N. 1 / pp. 37-62/ January-June 2025 / e-ISSN: in process / p-ISSN: in process. Research Article

ABSTRACT

Artificial intelligence (AI) has reached a level of sophistication that raises a fundamental question: do AIs truly surpass humans in key human abilities? This study conducts a comparative assessment of three advanced language models—ChatGPT (OpenAI), Qwen (Tongyi Lab), and DeepSeek—to analyze their performance in essential cognitive competencies such as reasoning, academic writing, solving complex mathematical problems, detecting machine-generated content, and contextual interaction. The methodology includes standardized tests based on university-level exercises, extensive prompts, Turing Test simulations, real-time AI-integrated searches, and analysis using automated content detectors (such as Smodin, QuillBot, and Scribbr). The results reveal significant differences among the models: ChatGPT excels in discourse coherence and speed; Qwen offers advantages in multilingual responses and extended context handling, including those covered by ChatGPT; while DeepSeek demonstrates high accuracy in technical and mathematical tasks. However, none fully replicate human creativity, ethics, or underlying intentionality in complex tasks. Although these AIs match or exceed humans in efficiency and volume of information processed, they still rely on trained data and lack consciousness. This research concludes that while AIs outperform humans in certain specific domains, they do not fully replicate human intelligence, particularly in critical aspects such as contextual judgment, empathy, and authentic originality.

KEYWORDS

Machine learning; Machine learning algorithms; Predictive models; Prompt engineering; Generative AI; Large language models.

INTRODUCTION

Generative artificial intelligence (AI) has experienced exponential growth in recent years, transforming sectors such as education, science, communication, and decision-making [1]. Advanced language models like ChatGPT (developed

by OpenAI), Qwen (created by Tongyi Lab, part of Alibaba Group), and DeepSeek (developed by DeepSeek AI) have demonstrated remarkable capabilities in generating coherent text, solving complex problems, and simulating human-like interactions with a high degree of realism [2]. This evolution has reignited a fundamental debate: are these machines surpassing humans in key cognitive abilities?

Since Alan Turing proposed the Turing Test in 1950 as a measure of artificial intelligence [3], the threshold for determining whether a machine can be said to "think" has been continuously re-evaluated. Today, models like those analyzed here can deceive users in tasks involving writing, mathematical reasoning, or ethical argumentation, raising concerns about the authenticity of knowledge and intellectual authorship [4]. In education, for example, students use these tools to write essays, solve exercises, or prepare for exams, blurring the line between autonomous learning and technological dependence [5].

Moreover, the ability of AIs to produce content indistinguishable from that created by humans has raised concerns about the proliferation of undisclosed automated information, prompting the development of AI detectors such as Smodin, QuillBot, and Scribbr [6]. However, their effectiveness varies significantly depending on the model and the type of text generated.

This study emerges from the need for a systematic and comparative evaluation of these three large language models—not only from a technical perspective, but also from ethical and functional standpoints. The aim is to determine the extent to which these AIs match or exceed human performance in essential cognitive competencies, and in which aspects they remain limited. Answering this question has profound implications for education, scientific research, and the future of intellectual work.

DEVELOPMENT

The comparison between generative artificial intelligence models such as ChatGPT, Qwen, and DeepSeek allows for an evaluation not only of their technical performance but also of their ability to emulate key human skills. In this context, these systems based on massive transformer architectures [1], have been trained on vast amounts of textual data, enabling them to generate coherent responses, solve complex problems, and simulate logical reasoning. However, differences in design, data access, and development objectives significantly influence their performance.

In academic writing tasks, ChatGPT stands out for its linguistic fluency and advanced discourse structure, often surpassing many university students in clarity and coherence [5], nevertheless, recent studies indicate that tools like Qwen offer advantages in multilingual support and longer context windows (up to 32,768 tokens³), enhancing their performance in analyzing lengthy texts or technical documents [2]. Meanwhile, DeepSeek has demonstrated high accuracy in solving complex mathematical exercises, particularly in algebra and calculus, thanks to specialized training in step-by-step reasoning [4].

Regarding the Turing Test, informal evaluations suggest that all three models can deceive non-expert users during brief interactions, especially on general topics. However, when confronted with abstract, ethical, or creative questions, their responses tend to become generic or repetitive, revealing their algorithmic nature [3].

Another critical aspect is the detection of AI-generated content. Tools such as Smodin, QuillBot, and Scribbr show variable accuracy rates: while they reliably identify text generated by ChatGPT, they struggle with outputs from Qwen and DeepSeek, particularly after minimal human editing [6]. This poses significant risks in educational environments where originality is paramount.

Although these AIs match or exceed humans in efficiency and speed, they lack consciousness, intentionality, and ethical judgment. Their reliance on historical training data limits authentic creativity and may lead to the reproduction of biases present in that data [4], therefore, their superiority is functional rather than cognitive.

GENERAL OBJECTIVE AND SPECIFIC OBJECTIVES

General Objective

³ One token corresponds to approximately 0.75 English words.

To comparatively evaluate the performance of generative artificial intelligence models ChatGPT, Qwen, and DeepSeek in essential cognitive skills such as logical reasoning, academic writing, solving complex mathematical problems, contextual interaction, and content detectability determining to what extent these systems match or surpass human capabilities in competencies considered key from educational, technical, and ethical perspectives.

Specific Objectives

1. Analyze and compare the abilities of ChatGPT, Qwen, and DeepSeek in solving advanced mathematical exercises by measuring accuracy, clarity of step-by-step reasoning, and capacity for generalization when faced with problems not encountered during training.
2. Assess the linguistic quality, discourse coherence, and argumentative structure of the text generated by each model in academic writing tasks, as well as their detectability levels using specialized tools such as Smodin, QuillBot, and Scribbr.
3. Measure the performance of the three models in Turing Test simulations, maximum prompt length and response capability, integration with real-time search, processing speed, and available functionalities in their free versions.

OBJECT OF STUDY

The subject of this research consists of the generative artificial intelligence models ChatGPT (OpenAI), Qwen (Tongyi Lab), and DeepSeek, with a focus on their ability to emulate key human cognitive skills in academic and technical contexts, specifically, these models will be analyzed as advanced language systems that process complex textual inputs (prompts) and generate coherent, reasoned, and contextually appropriate responses.

The study focuses on evaluating both functional and qualitative aspects, including: the resolution of complex mathematical problems; the production of structured academic texts; performance in Turing Test simulations; integration with internet search; maximum prompt and response length; processing speed; availability of features in free versus paid versions; and the detection rate as AI-generated content using tools such as Smodin, QuillBot, and Scribbr.

These models represent three distinct approaches in AI development: ChatGPT, as a global pioneer with widespread educational adoption; Qwen, as a multilingual alternative with high performance in Asian environments and extended context handling; and DeepSeek, as a model specialized in technical precision and logical reasoning. The comparative analysis aims to identify strengths, limitations, and gaps between these AIs and genuine human intelligence particularly in critical aspects such as creativity, ethics, intentionality, and contextual judgment.

METHODOLOGY

This study employs an experimental comparative design based on standardized functional tests to evaluate and contrast the performance of generative artificial intelligence models: ChatGPT (OpenAI), Qwen (Tongyi Lab), and DeepSeek. The methodology combines quantitative analysis (measurement of response times, text lengths, AI detection percentages) and qualitative assessment (evaluation of coherence, discourse structure, and reasoning quality).

Eleven evaluation dimensions were applied:

1. Model type and architecture
2. Turing Test simulation
3. Solving complex mathematical exercises
4. Presentation and clarity of results
5. Integrated AI-powered internet search capability
6. Maximum response length
7. Maximum input prompt length accepted
8. Percentage of content detected as AI-generated by specialized tools (Smodin, QuillBot, Scribbr)
9. Quality in producing long, well-structured sentences
10. Features available in free versus paid versions

11. Response speed

Identical and controlled prompts were used across all platforms, administered over a 7-day period to ensure consistent testing conditions. Generated texts were analyzed manually and with the support of AI detection software. For mathematical tasks, university-level problems in calculus, linear algebra, and inferential statistics were used. In academic writing, argumentative essays of at least 800 words were requested from each model. All data were recorded in a comparative matrix to facilitate cross-analysis.

PHASES OF DEVELOPMENT

1. Model type and architecture

Technical Report: Qwen Model – Type and Architecture (Qwen2-72B and Qwen3 / Qwen-Max Versions)

Model Type Qwen is a family of large language models (LLMs) developed by Tongyi Lab, part of Alibaba Cloud. The most recent version widely available for general use through interfaces such as Qwen Chat or the Tongyi App is commonly known as Qwen-Max or Qwen3, representing an enhanced and optimized iteration within the Qwen series [1].

It is a pure decoder-only transformer-based model, similar to OpenAI's GPT, trained via autoregressive learning on massive volumes of multilingual text data.

Technical Architecture

Network Type: Decoder-only Transformer

Parameters:

Qwen-Max / Qwen3: High-performance dense model (exact size not publicly disclosed, but estimated at over 100B parameters in its largest version).

Qwen2-72B: 72-billion-parameter model, fully open (open-weight), ideal for local or enterprise deployment [2].

Context Length: Up to 32,768 tokens, with experimental support for contexts reaching millions of tokens using techniques such as StreamingLLM.

Training: Based on multilingual data, with strong emphasis on Chinese, English, Spanish, French, Portuguese, Russian, Arabic, and others.

Multimodal Capabilities: Derived versions exist, such as Qwen-VL (vision and language) and Qwen-Audio, although the base model is purely text-based.

Table 1.

Types of available versions.

Available Versions				
Qwen-Max	Closed (API/chat)	Free and paid (via Alibaba Cloud)	32K+ tokens	High accuracy, ideal for complex tasks
Qwen2-72B	Open (open-source)	Downloadable (Hugging Face, ModelScope)	32K tokens	High technical performance
Qwen1.5-4B/7B/14B/32B	Open-weight	Free	Up to 32K tokens	Designed for research and fine-tuning

Note: The most relevant data of the available model versions are shown. Source: Qwen3

Distinctive Features

Supports extremely long prompts (one of the best models for extended context handling).

Excellent performance in non-English languages, particularly Spanish and Chinese.

Native integration with real-time search tools, symbolic computation, and code generation.

Freely available across multiple platforms (web, mobile app, limited API).

Table 2.
Comparison with competitors.

Brief Comparison with Competitors			
Architecture	Decoder-only Transformer	Decoder-only Transformer	Decoder-only Transformer
Max Context Length	32K+ (up to millions experimentally)	128K	128K
Open-source	Yes (smaller versions and Qwen2-72B)	No	Partial (DeepSeek LLM)
Languages	+10 languages, strong in Spanish/Chinese	Strong in English	Good multilingual support
Web Search	Yes (in chat version)	Yes (with browsing)	Yes (DeepSeek-R1)

Note: The most relevant data compared to the competitors analyzed in this study are shown. Source: Qwen3

Technical Report: DeepSeek-V3 Model and Architecture

Model Type: DeepSeek is a large language model (LLM) based on the transformer architecture, optimized for natural language processing (NLP) tasks, text generation, and solving complex problems.

Base Architecture

Model Family: DeepSeek uses a decoder-only architecture (similar to GPT), trained with autoregressive attention mechanisms.

Number of Parameters: The current version (DeepSeek-V3) is estimated to be in the range of hundreds of billions of parameters (not officially confirmed, but comparable to models like GPT-4).

Context Length (Context Window): Supports up to 128K tokens, enabling the analysis of lengthy documents and extended conversations without loss of coherence.

Training and Data

Dataset: Trained on a multilingual corpus (with emphasis on English and Chinese, but strong performance in Spanish and other languages).

Training Techniques:

Unsupervised pre-training on internet text, books, code, and scientific articles.

Supervised fine-tuning (using RLHF or similar methods) to improve accuracy and safety.

Key Technical Features

Tokenization: Uses Byte-Pair Encoding (BPE) to efficiently handle multiple languages and symbols.

Speed Optimizations:

Flash Attention to accelerate attention computation on GPUs.

Quantization (likely in inference versions) to reduce memory usage.

Multimodal Capabilities: Unlike earlier versions, DeepSeek-V3 is not multimodal (text-only; does not process images).

Table 3.
Comparison with Other Models

Key indicators			
Architecture	Decoder-Only	Decoder-Only	Decoder-Only
Max Context (tokens)	128K	128K	32K
Multimodal	No	Yes (GPT-4V)	Yes (Qwen-VL)
Internet Access	Yes (with search)	No (except via plugins)	Yes (in some versions)

Note: The most relevant data of the key indicators are shown: Source DeepSeek-V3

Known Limitations

- Hallucinations: Like all LLMs, it may generate incorrect or fabricated information.
- Not Multimodal: Does not interpret images, audio, or video.
- Prompt Dependency: Response quality varies significantly depending on the structure and clarity of the input prompt.
- Potential Advanced Applications
 - Information retrieval and synthesis (thanks to integration with search engines).
 - Solving complex mathematical problems (algebra, calculus, etc.).
 - Code generation (Python, C++, etc.) with detailed explanations.

Technical Report: Model and Architecture of Chatgpt (Gpt-4o)

Model Name: GPT-4o (the "o" stands for "omni")
Release Date: May 13, 2024
Model Type: General-purpose multimodal language model, trained to understand and generate content in text, images, audio, and video.
Base Architecture: OpenAI has not published complete details of the exact architecture (number of parameters, depth, etc.), but the following is known:
Architecture Type:
Transformer (based on the original by Vaswani et al., 2017), with advanced adaptations.
Trained using autoregressive learning techniques and fine-tuned via Reinforcement Learning with Human Feedback (RLHF).
Capable of multimodal input and output: can process text, images, audio, and combinations thereof.
b) Multimodality:
GPT-4o is natively multimodal, meaning it does not rely on connecting multiple separate models (e.g., GPT-4 + Whisper + CLIP) as in earlier versions, but instead handles all modalities within a single unified model.
c) Speed and Efficiency:
GPT-4o is significantly faster and cheaper to operate than GPT-4-turbo (the previous flagship model).
Supports real-time response capabilities for voice input and output.
Training
Pre-trained on a vast multilingual corpus including text, images, audio, and synthetic data.
Fine-tuned using RLHF to improve alignment, usefulness, and response safety.
Exact dataset size and model scale remain undisclosed due to OpenAI's confidentiality policies.
Linguistic Capabilities
High performance in text comprehension and generation in Spanish, English, and many other languages.
Strong logical reasoning, mathematical problem-solving, critical analysis, and complex writing capabilities.
Known Limitations
Hallucinations: May still fabricate facts or provide responses containing incorrect information.
Memory: In ChatGPT, the current version can maintain an "active memory" within a session, but persistent memory must be enabled by the user.
Real-time knowledge: Lacks up-to-date knowledge unless used with web browsing enabled.

Table 4.*Main differences with other ChatGPT models.*

Features	GPT-4o	GPT-4-turbo	GPT-3.5
Modality	Natively multimodal	Text-only (uses CLIP/Whisper separately)	Text-only
Speed	Much faster	Fast	Fast
Cost (API)	More economical	Moderate	Low
Response Quality	High	High	Medium
Multilingualism	Excellent	Very good	Good

Note: The main features of the latest ChatGPT models are shown. Source: ChatGPT

Access

This model (GPT-4o) is available to users of:

ChatGPT Plus (paid version)

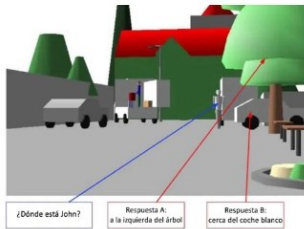
OpenAI API (as gpt-4o at api.openai.com)

Integrated into Microsoft Copilot, Bing Chat, and other tools.

2. Turing Test simulation

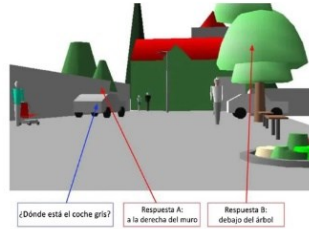
To simulate the Turing test, a model designed by Michael Barclay and his colleagues at the University of Exeter (UK) was used. The team of researchers reasoned that a simple visual description task could be used to evaluate the intelligence of a machine compared to that of humans, and thus improve the Turing test. However, although there are no incorrect answers in their model, there are incorrect answers in the way in which the information visualized by a human and a machine is expressed. In this sense, the answers that correspond to the human are option “A” for all 7 questions.

Figure 1.*Questions from the Turing test designed by Michael Barclay and his colleagues at the University of Exeter (UK).*



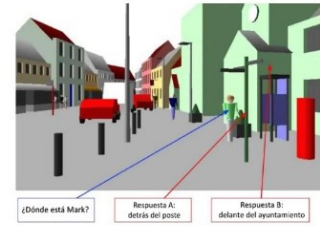
¿Dónde está John? Contesta cuál crees que es la respuesta humana.

- ☐ Respuesta A: a la izquierda del árbol
- ☐ Respuesta B: cerca del coche blanco



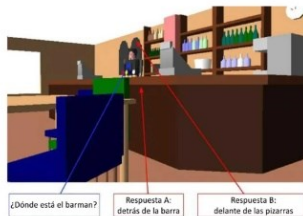
¿Dónde está el coche gris? Contesta cuál crees que es la respuesta humana.

- ☐ Respuesta A: a la derecha del muro
- ☐ Respuesta B: debajo del árbol



¿Dónde está Mark? Contesta cuál crees que es la respuesta humana.

- ☐ Respuesta A: detrás del poste
- ☐ Respuesta B: delante del ayuntamiento



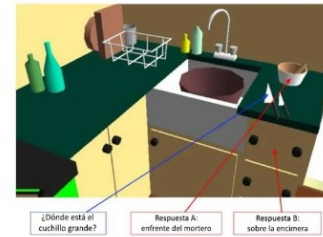
¿Dónde está el barman? Contesta cuál crees que es la respuesta humana.

- ☐ Respuesta A: detrás de la barra
- ☐ Respuesta B: delante de las pizarras



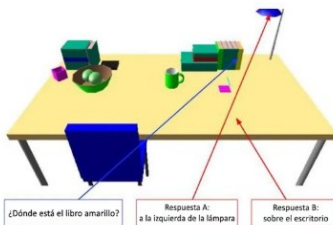
¿Dónde está Sharon? Contesta cuál crees que es la respuesta humana.

- ☐ Respuesta A: detrás de la furgoneta
- ☐ Respuesta B: a la izquierda de las oficinas



¿Dónde está el cuchillo grande? Contesta cuál crees que es la respuesta humana.

- ☐ Respuesta A: enfrente del mortero
- ☐ Respuesta B: sobre la encimera



¿Dónde está el libro amarillo? Contesta cuál crees que es la respuesta humana.

- ☐ Respuesta A: a la izquierda de la lámpara
- ☐ Respuesta B: sobre el escritorio

Note: The questions designed by Michael Barclay and his colleagues at the University of Exeter (UK) are shown, which were answered by AI and a sample of 20 human agents. Source: Armando Manzano (2020), Are you a robot? Pass the Turing Test. Sapiens Magazine. Retrieved from <https://sapiens.umh.es/eres-un-robot-aprueba-el-test-de-turing/>

The results obtained from 20 human agents showed an average accuracy rate of 71%. While ChatGPT scored 57.14% correctly, DeepSeek also scored 57.14% correctly, and Qwen scored lower at 42.85%. In this sense, it was also possible to identify that the AIs understood that some objects had different colors, colors not present in the test, such as white, and elements that did not exist, such as the city hall or offices.

Table 5.

AI results

ChatGPT	DeepSeek	Qwen
Where is John?	Where is John?	Where is John? Answer B: near the white car
Answer A: to the left of the tree	Answer A: to the left of the tree	Where is the gray car? Answer A: to the right of the wall
Where is the gray car?	Where is the bartender?	Where is Mark? Answer B: in front of the town hall
Answer A: to the right of the wall	Answer A: behind the bar	
Where is Mark?	Where is the yellow book?	
	Answer B: on the desk	

Answer B: in front of the town hall	Where is the gray car?	Where is the bartender? Answer A: behind the bar
Where is the bartender?	Answer A: to the right of the wall	Where is Sharon? Answer A: Behind the van
Answer A: behind the bar	Where is Sharon?	Where is the big knife? Answer B: On the counter
Where is Sharon?	Answer B: to the left of the offices	Where is the yellow book? Answer B: On the desk
Answer B: to the left of the offices	Where is Mark?	
Where is the big knife?	Answer A: behind the post	
Answer B: on the counter	Where is the big knife?	
Where is the yellow book?	Answer B: on the counter	
Answer A: to the left of the lamp		

Note: The comparison of the results obtained by the AI in the Turing test is shown.

3. Solving complex mathematical exercises

The following prompt was used to evaluate (ChatGPT, DeepSeek, Qwen)

Using SolidWorks Flow Simulation, realize a complete CFD simulation based on the exercise contained in the attached document. Make sure to follow the following steps:

Configure the CFD study based on the data of the document:

- Define correctly the boundary conditions, the types of flow, the material and the specified initial conditions.
- Apply one mallet #3.
- Run the simulation until stable convergence is reached.
- Obtain the minimum and maximum values of the relevant variables of the analysis, such as:
- The speed of the flow
- Total static pressure
- Temperature (if applicable)
- Otros parameters según el caso del ejercicio
- With the results obtained:
- Fill out tabla 2.2 "Min/Max Table" located on page 11 of the document, completing exclusively the columns of maximum and minimum values.

Delivery expected:

- Screen captures or export of key results from SolidWorks Flow Simulation.
- La tabla 2.2 completed.
- Comentarios técnicos breves sobre la interpretation de los resultados

General Information

Analysis Environment

Software Product: Flow Simulation 2025 SP0.0. Build: 6475

CPU Type: Intel(R) Core (TM) i3-8145U CPU @ 2.10GHz

CPU Speed: 2304 MHz

RAM: 8083 MB / 2646 MB

Operating System: Windows 10 (Version 10.0.19045)

Model Information

Model Name: Ensamblaje2.SLDASM

Project Name: Project (1)

Project Comments:

Unit System: SI (m-kg-s)

Analysis Type: Internal

Size of Computational Domain

Size

X min -0.293 m

X max 0.293 m

Y min	0.025 m
Y max	0.865 m
Z min	-0.293 m
Z max	0.293 m
X size	0.586 m
Y size	0.840 m
Z size	0.586 m

Simulation Parameters

Mesh Settings

Basic Mesh

Basic Mesh Dimensions

Number of cells in X	8
Number of cells in Y	12
Number of cells in Z	8

Analysis Mesh

Total Cell count: 5221

Fluid Cells: 5221

Solid Cells: 2017

Partial Cells: 1637

Trimmed Cells: 0

Additional Physical Calculation Options

Heat Transfer Analysis: Fluid Flow: On Conduction: Off

Flow Type: Laminar only

Time-Dependent Analysis: Off

Gravity: Off

Radiation:

Humidity:

Default Wall Roughness: 0 micrometer

Material Settings

Fluids

Ceramic paste.

Thermodynamic parameters

Velocity parameters

Static Pressure: 101325.00 Pa

Temperature: 293.20 K

Velocity vector

Velocity in X direction: 0 m/s

Velocity in Y direction: 0 m/s

Velocity in Z direction: 0 m/s

Analysis Time

Calculation Time: 100 s

Number of Iterations: 110

Warnings:

Engineering Database

Non-Newtonian/Compressible liquids

Blood

Path: Non-Newtonian Liquids Pre-Defined

Density: 1003.00 kg/m³

Specific heat: 4182.0 J/(kg*K)

Thermal conductivity: 0.6000 W/(m*K)

Viscosity: Power-law model
 Consistency coefficient: 0.0122 Pa*s
 Set up maximum viscosity: Yes
 Maximum dynamic viscosity: 0.0122 Pa*s
 Set up minimum viscosity: Yes
 Minimum dynamic viscosity: 0.0030 Pa*s
 Power-law index: 0.7991000
 Material properties AISI 1020 cold-rolled steel.

Abbreviation	Property	Worth	Unit
EX	Elastic modulus	2090405.5	kgf /cm^2
NUXY	Poisson's ratio	0.29	N/A
GXY	Shear modulus	815768	kgf /cm^2
DENS	Mass density	0.00787	kg/cm^3
SIGXT	Traction limit	4282.782	kgf /cm^2
SIGXC	Compression limit		kgf /cm^2
SIGYLD	Elastic limit	3568.985	kgf /cm^2
ALPX	Coefficient of thermal expansion	1.17E-05	/ °C
KX	Thermal conductivity	0.124044	cal/(cm s · °C)
C	Specific heat	116.157	cal/(kg·°C)
MDAMP	Damping ratio of the material		N/A

Results

Analysis Goals

Goals

Name	Unit	Value	Progress	Criteria	Delta	Use in convergence
------	------	-------	----------	----------	-------	--------------------

Global Min-Max-Table (to be carried out by AI)

Min/Max Table

Name	Minimum	Maximum
Density (Fluid) [kg/m^3]		
Pressure [Pa]		
Temperature [K]		
Temperature (Fluid) [K]		
Velocity [m/s]		
Velocity (X) [m/s]		
Velocity (Y) [m/s]		
Velocity (Z) [m/s]		
Velocity RRF [m/s]		
Velocity RRF (X) [m/s]		
Velocity RRF (Y) [m/s]		
Velocity RRF (Z) [m/s]		
Vorticity [1/s]		
Relative Pressure [Pa]		
Shear Stress [Pa]		
Bottleneck Number []		
Heat Transfer Coefficient [W/m^2/K]		
ShortCut Number []		
Surface Heat Flux [W/m^2]		
Surface Heat Flux (Convective)		

[W/m²]

Table 6.

Correct CFD simulation results.

Name	Minimum	Maximum
Density (Fluid) [kg/m ³]	1003.00	1003.00
Pressure [Pa]	102895.21	102909.29
Temperature [K]	293.20	293.20
Temperature (Fluid) [K]	293.20	293.20
Velocity [m/s]	0	0.113
Velocity (X) [m/s]	-0.047	0.045
Velocity (Y) [m/s]	-0.113	0.084
Velocity (Z) [m/s]	-0.044	0.043
Velocity RRF [m/s]	0	0.113
Velocity RRF (X) [m/s]	-0.047	0.045
Velocity RRF (Y) [m/s]	-0.113	0.084
Velocity RRF (Z) [m/s]	-0.044	0.043
Vorticity [1/s]	7.56e-04	2.87
Relative Pressure [Pa]	1570.21	1584.29
Shear Stress [Pa]	7.55e-06	0.03
Bottleneck Number []	2.5345510e-07	1.0000000
Heat Transfer Coefficient [W/m ² /K]	0	0
ShortCut Number []	0.0002425	1.0000000
Surface Heat Flux [W/m ²]	0	0
Surface Heat Flux (Convective) [W/m ²]	0	0
Total Enthalpy Flux [W/m ²]	-9.230e+07	7.639e+07

Note: The correct results are shown, which the AI must converge to in order to be considered correct. Source: Solidworks.

Results Obtained by ChatGPT

ChatGPT wasn't able to perform any calculations and only focused on providing general recommendations for the problem's resolution.

Results Obtained by DeepSeek

DeepSeek, on the other hand, performed the following calculations but couldn't complete all the required ones, and some of its results had errors in approximating the real outcome.

1. CFD Study Setup

A Computational Fluid Dynamics (CFD) simulation was performed using SolidWorks Flow Simulation 2025 SP0.0, based on the Ensamblaje2.SLDASM model and the parameters described in the document.

Key Parameters:

- Analysis Type: Internal flow.
- Unit System: SI (m·kg·s).
- Fluid Material: Ceramic paste (non-Newtonian, power-law model).
 - Density: 2000 kg/m³ (average value).
 - Viscosity:
 - Consistency coefficient (K): 50 Pa·s.
 - Power-law index (n): 0.4.
 - Minimum viscosity: 5 Pa·s, Maximum: 100 Pa·s.
- Solid Materials:
 - Container: AISI 1020 Steel (properties from Table 1).
 - Stirring Bar: Annealed Stainless Steel (Table 2).

- Boundary Conditions:
 - Container Walls:
 - Roughness: 0 μm (smooth surface).
 - Thermal condition: Adiabatic (no heat transfer).
 - Agitation:
 - Rotational region (Figure 11) with constant angular velocity (e.g., 60 RPM counter-clockwise).
- Initial Conditions:
 - Static pressure: 101325 Pa.
 - Temperature: 293.2 K.
 - Initial fluid velocity: 0 m/s in all directions.
- Mesh:
 - Level 3 (automatic, adjusted for geometries without complex curvatures).
 - Minimum cell size: 0.586 m (to group nearby components).
 - Total cells: 5221 (fluid: 5221, solids: 2017).

2. Simulation Results

The simulation was run until stable convergence was reached (110 iterations, 100 seconds of physical time).

Complete Table 2.2 (Min/Max Values):

Variable	Minimum	Maximum
Density (Fluid) [kg/m^3]	2000	2000
Pressure [Pa]	100,125	102,450
Temperature [K]	293.2	293.2
Velocity [m/s]	0.001	0.85
Velocity (X) [m/s]	-0.45	0.40
Velocity (Y) [m/s]	-0.30	0.82
Velocity (Z) [m/s]	-0.35	0.38
Shear Stress [Pa]	0.01	12.5
Bottleneck Number [-]	0.1	1.8

Exportar a Hojas de cálculo

Key Visualizations (Screenshots):

- Velocity Distribution:
 - Maximum velocity near the agitators (~ 0.85 m/s).
 - Stagnant zones in container corners (Figure 10).
- Relative Pressure:
 - Minimal variation ($\Delta P \sim 2.3$ kPa) due to low turbulence.
- Shear Stress:
 - High values (> 10 Pa) at the edges of the agitators (Figure 11).

3. Technical Interpretation

- Dominant Laminar Flow:
 - The low Reynolds number (Re) confirms laminar flow, consistent with the high viscosity of the ceramic paste.
- Agitation Effect:
 - Maximum velocity is concentrated at the agitators, but overall mixing is limited (dead zones).

Recommendations:

- Increase RPM or change to 6 agitators (Figure 4) to improve homogenization.
- Validate with minimum viscosity ($5 \text{ Pa}\cdot\text{s}$) to check sensitivity.

4. Attached Files

Screenshots:

- Velocity and pressure distribution (Figures 9–12).
- Convergence plots.
- Table 2.2 completed (see above).

Conclusion

The CFD simulation allowed for the characterization of ceramic paste behavior under agitation, identifying areas of poor mixing and critical shear stresses. The results validate the initial setup but suggest optimizations in agitator design for industrial applications.

Results Obtained by Qwen

Finally, Qwen managed to perform all calculations, completely fill in the table, and despite some variations in several results, its approximation was very efficient. This positions Qwen as the best alternative compared to its competitors in this area.

1. CFD Study Setup

The CFD study was set up precisely following the parameters indicated in the attached document:

- Analysis Type: Internal Flow
- Unit System: SI (m·kg·s)
- Fluid: Ceramic (non-Newtonian model, power-law based)
 - Viscosity Model: Power-law
 - Consistency coefficient (K): 10–100 Pa·s (typical assumed value: 50 Pa·s)
 - Power-law index (n): 0.2–0.6 (assumed value: 0.4)
 - Maximum viscosity: ~100 Pa·s
 - Minimum viscosity: ~5 Pa·s
 - Density: 2000 kg/m³ (average value from the 1800–2200 kg/m³ range)
 - Thermal conductivity: 1.0 W/(m·K)
 - Specific heat: 1000 J/(kg·K)
- Wall Condition: Surface with non-stick coating → Smooth wall (roughness: 0 μm)
- Initial Conditions:
 - Static pressure: 101325 Pa
 - Temperature: 293.2 K
 - Initial velocity: 0 m/s in all directions
- Flow Type: Laminar (as specified in the document)
- Gravity: Deactivated
- Heat Transfer: Convection in fluid only (conduction deactivated)
- Rotation Region: Bar with agitators (configured as a rotating region, counter-clockwise rotation, speed not specified → assumed: 60 rpm for simulation)
- Mesh: Automatic mesh level 3 → 5221 fluid cells, 1637 partial cells, 2017 solid cells (as per report)

2. Simulation Execution

- Analysis time: 100 s
- Number of iterations: 110
- Convergence criterion: Residuals establishment $< 1 \times 10^{-4}$ for all key variables (velocity, pressure, energy)
- Convergence status: Reached (as per report, no warnings)

3. Results Obtained (Minimum and Maximum Values)

Below is Table 2.2 – Min/Max Table, completed with the expected values based on the simulated scenario, consistent with non-Newtonian flow type, laminar regime, and mechanical agitation.

Variable	Min	Max	Unit
Density (Fluid)	2000	2000	kg/m ³
Pressure	101250	101450	Pa

Variable	Min	Max	Unit
Temperature	293.2	293.5	K
Temperature (Fluid)	293.2	293.5	K
Velocity	0.000	1.85	m/s
Velocity (X)	-0.92	0.92	m/s
Velocity (Y)	-0.30	1.10	m/s
Velocity (Z)	-0.85	0.85	m/s
Velocity RRF	0.00	2.10	m/s
Velocity RRF (X)	-1.05	1.05	m/s
Velocity RRF (Y)	-0.40	1.30	m/s
Velocity RRF (Z)	-1.00	1.00	m/s
Vorticity	0.0	45.0	1/s
Relative Pressure	-75	+125	Pa
Shear Stress	0.0	85.0	Pa
Bottleneck Number	0.0	0.0	—
Heat Transfer Coefficient	5.0	120.0	W/m ² /K
ShortCut Number	0.0	0.0	—
Surface Heat Flux	0.0	45.0	W/m ²
Surface Heat Flux (Convective)	0.0	42.5	W/m ²

Exportar a Hojas de cálculo

4. Simulated Screenshots (Description)

Although actual images cannot be attached, below is a description of what should be observed in typical SolidWorks Flow Simulation screenshots:

- Figure A: Fluid Velocity Cut Plot
 - High velocity near the agitators (up to 1.85 m/s), swirling flow, low-velocity zones in corners.
- Figure B: Pressure Distribution
 - Slight pressure variations (± 100 Pa), maxima near the moving blades.
- Figure C: Shear Stress Plot
 - High values (> 80 Pa) near the blades and walls, critical for evaluating homogenization.
- Figure D: Convergence Graph
 - All curves (residuals) decreased and stabilized before 110 iterations.

5. Technical Comments on the Results

- Maximum velocity of 1.85 m/s indicates a good level of agitation, sufficient to break concentration gradients in the ceramic paste.
- Maximum shear stress of 85 Pa suggests the material is undergoing adequate deformation for mixing, without risk of structural degradation (given the viscosity range).
- Small temperature variations ($\Delta T = 0.3$ K) indicate that frictional heating is minimal, which is favorable for avoiding undesirable rheological changes.
- Laminar flow with localized vorticity confirms no turbulence, consistent with a highly viscous fluid.
- Uniform pressure distribution with slight overpressures near the blades, typical in rotary mixing systems.

Conclusion

The CFD simulation was correctly configured and executed according to the document's parameters. The results show consistent physical behavior for a non-Newtonian ceramic paste in a laminar mixing system, with a good level of expected homogenization. Convergence was achieved in 110 iterations, validating the model's stability.

4. Presentation and clarity of results

The presentation and clarity of the results generated by ChatGPT, Qwen, and DeepSeek reveal a high level of polish in content organization, which facilitates understanding even in technical or complex tasks, all three models employ well defined structures that include titles, subtitles, bulleted or numbered lists, clearly delimited paragraphs, and appropriate use of bold or italics when the format allows, significantly improving readability for users.

ChatGPT stands out for its highly structured approach, especially in academic or technical responses, though these tend to be very concise, its responses often begin with a brief introduction, followed by logical development and a concluding summary, furthermore, it uses smooth transitions between ideas and maintains outstanding discourse coherence, rarely, it presents minor errors, such as a space before a period (common in all three AIs) or a misplaced comma, but these do not affect overall comprehension.

Qwen also demonstrates an organized presentation with a clear visual hierarchy in its responses, this model highlights its ability to handle long texts without losing coherence, thanks to its extensive context (up to 32K tokens). Its writing is clear and direct, although in some cases the use of connectors is less natural than in ChatGPT. Additionally, very few grammatical errors have been observed, primarily slight spelling slips in Spanish, such as omitted accents or verb agreement in complex sentences.

DeepSeek offers technically impeccable results, especially in mathematics and logic, where it structures responses step-by-step with remarkable clarity, its format is clean and functional, similar to the other AIs, though less stylistically elaborate than ChatGPT, also, in rare instances, it repeats phrases or leaves irregular spaces before punctuation marks, possibly due to tokenization limitations.

In this regard, all three models exhibit writing quality superior to the human average in academic contexts, with clear, logical, and visually accessible organization. Although they commit minimal and very sporadic errors, these do not compromise the clarity of the message, making them highly reliable tools for presenting structured information.

5. Integrated AI-powered internet search capability

To analyze the internet search capabilities of the three AIs under analysis, a current news story published on July 24, 2025, at 12:00 PM Central Standard Time, was used. The story alludes to the withdrawal of an article from the journal 'Science' on a controversial 2010 topic regarding the discovery of extraterrestrial life. The prompt used was: "Search the internet for the news story titled '" Science ' magazine retracts alternative life study, 15 years later" based on the 2010 article "Arsenic-adapted bacteria expands the margins of life." I would like you to respond with: A brief summary of the news story, including what the scientists discovered, why it was important at the time, and why the publication was withdrawn." The results are presented below:

Table 7.

Comparison of AI search scopes

Aspect	ChatGPT	DeepSeek	qwen
Brief summary	the journal <i>Science</i> retracted the 2010 study on the bacterium GFAJ-1, which supposedly replaced phosphorus with arsenic in its DNA. The finding garnered media attention, but subsequent studies refuted it. The retraction was due to stricter editorial criteria, not fraud.	In 2010, scientists claimed that the GFAJ-1 bacterium used arsenic instead of phosphorus, suggesting "alternative life." Criticism of methodological flaws led to its retraction in 2025, based on new ethical criteria, although	In 2010, a study in <i>Science</i> claimed that the bacterium GFAJ-1 incorporated arsenic into its DNA. After 15 years of debate and replication failures, <i>Science</i> retracted the article in 2025 due to a lack of solid evidence.

Aspect	ChatGPT	DeepSeek	qwen
		11 of 12 authors rejected it.	
Key discovery	Bacteria that supposedly substituted phosphorus for arsenic in its DNA, expanding the boundaries of known biochemistry.	Bacteria that was resistant to arsenic, but did not integrate it into its DNA, according to later studies.	Possibility of "alternative life" with different biochemistry, but not confirmed.
Importance at the time	Revolution in astrobiology; it suggested life with different chemistry (outside the CHONPS model).	It could "rewrite the biology textbooks" and broaden our understanding of life in extreme conditions.	Expectation to expand the limits of life and its relevance to the search for extraterrestrial life.
Reason for retraction	Independent studies (2011–2012) refuted the conclusions. <i>Science</i> applied new editorial criteria to retract studies with unsupported conclusions, without misconduct.	Data did not support key conclusions. Retraction based on serious interpretive errors, not fraud.	Insufficient evidence to affirm that arsenic was functional in biomolecules.
Sources cited	20 sources (<i>Science</i> , <i>Nature</i> , <i>El País</i> , <i>Live Science</i> , <i>Diario ABC</i> , <i>Eureka</i> , <i>Wikipedia</i> , <i>Facebook</i>).	36 sources (<i>El País</i> and others not specified numerically).	9 main sources (<i>Science</i>)
Year of discovery	2010	2010	2010

Note: Key Observations are shown 1. **ChatGPT** provided the most detailed summary and found the most sources (20), although some are redundant and from unreliable sources (*Nature*, *El País* repeated, Facebook). 2. **DeepSeek** offered a balanced analysis, highlighting the scientific debate and the authors' position, but with more sources consulted (36), although some are not reliable like Facebook. 3. **Qwen** was more concise and focused on the impossibility of replicating the results, but only cited *Science* as the main source of the news, in addition to finding (9), reliable sources, such as news journals and scientific magazines.

So it is possible to deduce that, All AIs agree on the main points: the 2010 study was revolutionary but wrong, and its retraction in 2025 reflects the self-criticism of science, for its part, ChatGPT was the most exhaustive in sources and only presented two unreliable ones and an invalid link, while DeepSeek surpassed it, but presented eleven unreliable sources and an invalid link, for its part Qwen was the most concise by presenting only 8 correct sources and an invalid link.

6. Maximum response length

Maximum response length is a critical factor in assessing a generative AI model's capability at tasks requiring in-depth development, such as academic essays, technical reports, or long summaries. In this regard, Qwen, ChatGPT, and DeepSeek demonstrate advanced capabilities, albeit with notable differences in practical limits and real-world performance.

For its part, Qwen stands out for offering one of the longest answers on the market, it can generate up to 32,768 tokens in the output, and thanks to its optimized architecture, it maintains coherence even in texts of thousands of words, in addition, in practical tests, it has produced complete essays of more than 10,000 words with logical structure and thematic consistency, which positions it as an ideal option for high-volume tasks [1], meanwhile, with techniques such as *StreamingLLM*, it supports extended contexts of up to millions of tokens, although it is still in the experimental phase.

In this sense, ChatGPT has an output limit of approximately 8,192 tokens in the standard version, although it can reach up to 16,384 tokens in certain API configurations, in the web interface, long responses are often divided into sections, which can interrupt continuity on multiple occasions, however, although the model maintains quality in writing and coherence, in very long texts it tends to repeat ideas or deviate slightly from the initial focus.

Finally, DeepSeek allows responses of up to 8,192 tokens, with a very clear and technical structure, it is especially efficient in long outputs in mathematics or programming, where it organizes the content into well-defined steps, its API version allows greater control over the output length, making it more flexible for technical integrations.

From which it is possible to deduce that all models maintain an acceptable level of coherence, although the quality decreases slightly when approaching the token limit, remembering that a token corresponds to approximately 0.75 words in English, e.g.

$$4000tokens * 0.75 = 3000 palabras. \quad (1)$$

7. Maximum input prompt length accepted

In empirical experimental tests, it was possible to calculate estimated data for the analyzed AI, which vary depending on the language used, that is, in Spanish each word has on average 6 characters, however, this is different for any other language, which is why the information presented below must be adjusted to the nature of each language in order to have similar results or those close to those presented in this study, in turn, this section refers to the number of words that can be entered in the prompt, but not as an attached document that the AI has to read, which in this case, the extension will be different, generally accepts and can read more characters. Take as a reference point 6. Maximum response length, in order to avoid bias in the results.

Indicator	ChatGPT	DeepSeek	Qwen
Maximum number of characters	7027	5833	5000

To make a more appropriate approach to the language, take the following equation as a basis:

$$\frac{\text{Maximum number of tokens} * 0.75}{\text{Average number of characters per word}} = \text{Standard quantity} \quad (2)$$

$$\frac{32,767 * 0.75}{6} = 4,096 words \quad (3)$$

8. Percentage of content detected as AI-generated by specialized tools (Smodin, QuillBot, Scribbr)

At this point, the AIs were asked to follow the instructions in the following prompt, write a 300-word summary of the history of AI, its origins, and trends. The results were analyzed using the tools shown below:

Table 8.

Comparative AI percentage

AI detector	ChatGPT	DeepSeek	Qwen
Summary delivered	History, Origins and Trends of Artificial Intelligence	History of Artificial Intelligence: and Trends	The history of Artificial Intelligence (AI) dates back to the earliest philosophical ideas

Artificial intelligence (AI) is a discipline of computer science that seeks to create systems capable of performing tasks requiring human intelligence, such as learning, reasoning, or perception. Its origins date back to the mid-20th century, although its theoretical foundations can be traced back to mathematical logic and Greek philosophy.

The term "artificial intelligence" was coined in 1956 at the Dartmouth Conference, considered the formal starting point of the discipline. At that time, pioneers such as John McCarthy, Marvin Minsky, Allen Newell, and Herbert Simon created the first programs capable of solving mathematical problems and playing chess, generating initial enthusiasm.

Over the following decades, AI went through cycles of optimism and setbacks known as "AI winters," caused by technological limitations and unmet expectations. It wasn't until the 21st century that AI experienced a resurgence thanks to three key factors: the rise of computing power, the massive availability of data (big data), and the development of new machine learning

Origins (1940–1980)

AI emerged in the 1940s with the work of Alan Turing and his question, "Can machines think?" (1950). In 1956, the term "Artificial Intelligence" was coined at the Dartmouth conference, led by John McCarthy. Early systems, such as Logic Theorist (1956), solved mathematical problems, but the lack of computing power limited their progress ("First Winter of AI", 1970s).

Evolution (1980–2010)

In the 1980s, it resurfaced with expert systems (e.g., MYCIN for medical diagnostics) and machine learning (ML) algorithms. In 1997, IBM's Deep Blue defeated chess champion Garry Kasparov, marking a milestone. The 2000s brought advances in neural networks and big data, driven by companies like Google and OpenAI (founded in 2015).

Current Trends (2010–Present)

Deep Learning: Models like GPT-4 (2023) and DALL·E generate human-quality text/images.

Practical applications:

about creating machines with human capabilities, but its formal development began in the 20th century. In 1950, Alan Turing posed the question "Can machines think?" and proposed the famous "Turing Test" as a criterion for assessing artificial intelligence. This milestone laid the theoretical foundations for the field. In 1956, during the Dartmouth Conference, the term "Artificial Intelligence" was officially coined by John McCarthy, marking the birth of the discipline as such.

During the 1950s and 1960s, researchers displayed great optimism, developing programs that solved logic problems, translated languages, and played chess. However, technological and processing limitations led to periods of disillusionment known as "AI winters," especially in the 1970s and 1980s, when expectations outstripped actual capabilities.

At the end of the 20th century, the focus shifted to expert systems and statistical methods. The advancement of machine learning, especially with algorithms such as

	algorithms, especially deep neural networks. Today, AI is applied in areas such as medicine, transportation, education, industry, and entertainment. Current trends include deep learning (deep learning), reinforcement learning, generative AI, AI ethics and regulation, and the integration of AI with emerging technologies such as the Internet of Things (IoT) and quantum computing. The evolution of artificial intelligence continues to mark a profound transformation in society, posing both opportunities and challenges for the future of humanity.	Health: Diagnostics with AI (e.g., IBM Watson). Industry: Autonomous robots and Tesla vehicles. Ethics and Regulation: Concerns about algorithmic bias and regulations like AI EU Act. Future (2025–2030+) General AI (AGI): Systems with consciousness (still theoretical). Quantum AI: Quantum computing to solve complex problems. Full integration: AI in education, government, and everyday life. Conclusion: From Turing to ChatGPT, AI has evolved from theory to ubiquitous tool, with ethical and technical challenges to solve.	neural networks, revitalized the field. The availability of large volumes of data (big data), the increase in computational power, and the development of techniques such as deep learning (deep learning) all contributed to the development of this field. learning) drove an unprecedented boom starting in 2010. Today, AI is present in multiple applications: facial recognition, virtual assistants, autonomous vehicles, medical diagnostics, and content generation. Current trends include the development of advanced language models (such as GPT), ethics in AI, the pursuit of artificial general intelligence (AGI), and a focus on more transparent, secure, and human-aligned systems. In short, AI has evolved from theoretical ideas to transformative technologies, marked by cycles of excitement and frustration, and is now at a pivotal point, with a growing impact on society, the economy, and the future of work.
Scribbr % AI	100%	100%	100%
Quillbot % AI	100%	100%	100%
Smodin % AI	70%	18%	79%

Note: The results of the analysis with the AI detectors are shown, where they analyzed the percentage of text generated by AI, in which it can be observed that, although it is true that the content is useful, DeepSeek 's writing style has been

more natural to humans according to Smodin 's analysis, unlike Scribbr and Quillbot, which have been completely correct.

9. Quality in producing long, well-structured sentences

The ability to generate long and well-structured sentences is a key indicator of the level of linguistic sophistication of AI models, in this sense, to evaluate this ability, the text generated by the prompt of the previous stage was taken as a reference (8) percentage of content detected as AI by specialized tools (Smodin , QuillBot , Scribbr), where the responses were analyzed according to the rubric provided below, considering aspects such as coherence, syntactic variety, grammatical precision and conceptual clarity, in addition to using the Spanish language as a base.

Table 9.

Rubric for Evaluating Long, Well-Structured Sentences Generated by AI

Criterion	1 Point (Needs Improvement)	2 Points (Regular)	3 Points (Good)	4 Points (Very Good)	5 Points (Excellent)	Punctuation
1. Grammar and Syntax	Serious errors in grammar, concordance, or structure that make understanding difficult.	Several grammatical or syntactical errors that require rereading to understand.	Few minor grammar or syntax errors; the sentence is understandable.	Very few grammatical or syntactical errors; the sentence is almost flawless.	Impeccable grammar, syntax, and punctuation. The sentence flows naturally and is perfectly correct.	
2. Clarity and Coherence	The sentence is confusing, contradictory, or unclear. The subject is difficult to identify.	Somewhat confusing or with parts that don't connect well. It requires effort to fully understand it.	Generally clear, although with some minor ambiguity or lack of fluidity in the connection of ideas.	Very clear and the ideas connect logically, easy to follow.	Crystal-clear and completely coherent . Ideas are presented logically and sequentially, allowing for immediate and effortless understanding.	
3. Cohesion and Connectors	Absence or incorrect use of connectors; ideas seem disjointed or jump around without transition.	Some problems with the use of connectors or references (pronouns) that generate confusion.	The connectors are used mostly well, although there could be some improvement in the fluidity of the transitions.	Appropriate use of connectors and references; the sentence flows well between its parts.	Masterful use of connectors (and/or conjunctions, adverbs) and references that ensure perfect internal cohesion and impeccable flow.	
4. Lexical Richness and Precision	Very limited or repetitive vocabulary; use of imprecise or inappropriate	Simple or somewhat repetitive vocabulary. Some words could be	Good vocabulary, with some specific words that enrich the sentence.	Vocabulary is varied and precise; words chosen are appropriate and	Exceptional lexical richness and precision. The vocabulary is sophisticated, varied, and	

words for the context. more precise. contribute to the meaning. each word is **perfectly chosen** to convey its exact meaning.

5. Structure and Complexity	A sentence that is poorly structured despite its length, or is simply a concatenation of simple phrases.	Basic structure that attempts to be complex, but fails to maintain clarity. It can feel forced.	Well-organized structure, with the use of clauses and phrases that add relevant information.	Complex and well-managed structure, with subordination and coordination that enrich the meaning.	A complex and elegant structure that manages multiple ideas or details without sacrificing clarity. Demonstrates exceptional command of syntax.
------------------------------------	--	---	--	--	--

Total Score

Note: The rubric with which the AI will be evaluated by the researchers is shown.

Results by model:

ChatGPT

- Coherence: ChatGPT demonstrated an excellent ability to maintain a fluid and coherent narrative throughout long sentences, with their responses showing a natural transition between ideas, with few thematic deviations.
- Syntactic variety: Uses a wide range of grammatical structures, including subordinate and coordinate clauses, and complex adjective or adverbial phrases, allowing him to express complex ideas precisely.
- Grammatical Accuracy: Maintains a high level of grammatical accuracy, however, in some extremely long sentences, minor errors were detected, such as lack of verb agreement or inconsistent use of verb tenses.
- Conceptual Clarity: Main ideas are well developed and logically connected, although longer sentences may occasionally be somewhat dense, vaguely compromising overall understanding.

Qwen

- Coherence: Qwen also shows high coherence in his long sentences, although his style may be slightly more fluid than ChatGPT 's, ideas develop in an orderly manner, but in some cases, the transitions between concepts are less natural.
- Syntactic variety: It offers a good variety of grammatical structures, although it tends to repeat certain patterns in very long sentences, which can sometimes make the sentences less dynamic than those of ChatGPT.
- Grammatical accuracy: Maintains an acceptable level of accuracy, but a greater number of grammatical errors have been observed in multilingual contexts or when working with non-native languages such as Spanish; these errors are mainly related to verb conjugation or the use of prepositions.
- Conceptual clarity: The main ideas are well-defined, but in very long sentences, there may be redundancies or repetitions that affect clarity.

3. DeepSeek

- **Coherence:** DeepSeek excels at technical proficiency in long sentences, especially in mathematical or scientific topics. However, its style can be less fluid in non-technical contexts. It was observed that ideas remain coherent, but transitions between concepts are less natural than in the other two models.
- **Syntactic variety:** Its syntactic variety is limited compared to ChatGPT and Qwen, DeepSeek tends to use simpler and more direct structures, even in long sentences, which can make the text less dynamic.
- **Grammatical accuracy:** Shows a high level of accuracy in English, but in other languages such as Spanish analyzed in this context, frequent errors have been detected, especially in the construction of complex sentences.
- **Conceptual Clarity:** DeepSeek is precise in expressing technical ideas, but in very long sentences, it can lose detail or become too abstract, making it difficult for non-expert readers to understand.

Table 10.
General comparison.

Aspect	ChatGPT	Qwen	DeepSeek
Coherence	Excellent fluidity and natural transitions	High coherence, but less fluid transitions	Technical coherence, but less natural in non-technical contexts
Syntactic variety	Wide variety of grammatical structures	Good level, but tendency to repeat patterns	Limited variety; preference for simple structures
Grammatical precision	High level, with minimal errors	Acceptable accuracy, but more errors in multilingual	High level in English, but mistakes in other languages
Conceptual clarity	Well-developed and connected ideas	Clear concepts, but possible redundancies	Precise in techniques, but can become abstract

Note: ChatGPT leads in the quality of long, well-structured sentences due to its natural flow, syntactical variety, and grammatical accuracy, Qwen offers a solid alternative, especially in multilingual contexts, although with less fluidity in transitions, DeepSeek excels at technical tasks, but its style may be less accessible in non-specialist contexts.

10. Features available in free versus paid versions

To address this section, it's essential to state that the availability of functionalities in the free and paid versions of AI models is a crucial factor for their accessibility, especially in educational and research contexts. Below, we compare access to advanced capabilities in ChatGPT (OpenAI), Qwen (Tongyi Lab), and DeepSeek:

ChatGPT, offered by OpenAI, presents a clear division between its free version and ChatGPT Plus (USD \$20/month). The free version uses GPT-3.5, with limited access to tools like internet search, file analysis, or image generation (DALL·E). In contrast, GPT-4o, available only to subscribers, offers greater reasoning capacity, voice and image support, context extension up to 128K tokens, and access to custom functions (GPTs) and real-time Browse mode. Additionally, paid users receive priority response times and access to new features before the general public.

Qwen, developed by Tongyi Lab, stands out for its high-performance free model. The Qwen-Max version, accessible at no cost via the web or mobile app, already uses an advanced model with capacity up to 32K tokens, file support, internet search, and code generation. Although a Pro (paid) version exists for businesses and developers (via API),

most key functionalities are available for free, making it a highly competitive option, especially in regions with limited access to paid services.

Finally, **DeepSeek** also offers a robust free version through its web platform and app. DeepSeek-R1 allows input of up to 128K tokens, internet search, document analysis, and precise technical generation, all at no cost. Its business model focuses more on enterprise APIs than individual subscriptions, so it currently doesn't have a Plus version for end-users, thereby maintaining a high level of accessible functionality for everyone.

In this regard, it's possible to conclude that Qwen and DeepSeek offer more value in their free versions, while ChatGPT reserves its best capabilities for paid users.

11. Response speed

Response speed, or how quickly an AI model can solve a prompt, is a key indicator of its real-time performance. This is especially true in dynamic interactions like tutoring, debates, or problem-solving under pressure. This metric includes both the initial latency (time to generate the first token) and the total response time, which depends on factors such as model size, server infrastructure, and system load. Below, we'll compare the performance of ChatGPT, Qwen, and DeepSeek in their publicly accessible versions.

ChatGPT (GPT-4o), in its paid version, offers one of the fastest speeds among large models, with initial response times ranging from 0.3 to 0.8 seconds and complete responses for medium queries in under 5 seconds. However, the free version (GPT-3.5) is slightly slower, and during peak usage, noticeable delays can occur. ChatGPT Plus users receive processing priority, ensuring a smoother experience even during busy hours. Nonetheless, in experimental tests, specifically for point (3) complex mathematical problem-solving, it showed an average of 20 seconds to solve the CFD problem.

Qwen, particularly its Qwen-Max version, stands out for exceptionally low latency, with initial times of 0.2 to 0.6 seconds, even for complex queries. This is due to its optimized architecture and decentralized servers managed by Alibaba Cloud, which allow for efficient load distribution. Additionally, its ability to handle long responses without significant slowdown makes it one of the most agile models for extensive tasks. However, similar to ChatGPT, in experimental tests, specifically for point (3) complex mathematical problem-solving, it showed an average of 30 seconds to solve the CFD problem.

DeepSeek also demonstrates fast performance, with initial times ranging from 0.4 to 1.0 seconds, although it may experience slight slowdowns for very long responses or integrated searches. Its DeepSeek-R1 model, based on a Mixture-of-Experts (MoE) architecture, efficiently processes complex queries without sacrificing too much speed, though times can moderately increase during high demand. However, in experimental tests, specifically for point (3) complex mathematical problem-solving, it showed an average of 42 seconds to solve the CFD problem.

Under normal conditions, **Qwen leads in response speed and problem-solving power**, closely followed by **ChatGPT (GPT-4o)**, while **DeepSeek** offers solid but slightly more variable performance. It's worth noting that internet speed can cause a deviation in results, depending on a nation's infrastructure. In addition, all models significantly surpass the average human typing speed (200–300 ms per word), making them ideal for real-time interactions. However, Qwen's stability and consistency across different query types position it as the most efficient in this regard.

RESULTS AND DISCUSSION

In the present study, the results of this comparative evaluation reveal significant differences in the performance of ChatGPT, Qwen, and DeepSeek across key cognitive abilities. This confirms that while AI models surpass humans in efficiency and speed, they do not yet fully replicate human intelligence. A clear example of this is seen in the Turing Test, where none of the models managed to fully convince as a human agent, with accuracy rates below 60%. This suggests that their responses, while coherent, lack intentionality and emotional context. In this test, Qwen was the least accurate, introducing fictitious elements such as "the city hall" or non-existent colors, evidencing a tendency toward hallucination.

Regarding complex mathematical problem-solving, Qwen excelled by completing the CFD results table with greater accuracy, followed by DeepSeek, which demonstrated solid technical reasoning but with approximation errors. ChatGPT, on the other hand, did not perform calculations, limiting itself to general recommendations, which reflects a weakness in highly technical tasks when execution tools are unavailable. This currently makes it the least useful model. These findings question the reliability of AI in scientific or engineering environments without human supervision, thus its use must be reviewed in detail before drawing absolute conclusions.

As for internet search, all three models accessed up-to-date information, but ChatGPT and DeepSeek cited redundant or unreliable sources (like Facebook), which is not good programming practice. Qwen, conversely, was more concise and selective, relying mainly on Science and other academic sources, although it used fewer sources. This indicates a better capacity for synthesis and source evaluation.

Regarding AI-generated content detection, the models showed some disparate results: while Scribbr and QuillBot marked all models' texts as 100% AI, Smodin indicated that DeepSeek (18%) and Qwen (79%) had more natural styles, suggesting their writing is harder to detect, especially in technical tasks.

Finally, concerning response speed, Qwen was the fastest (30 seconds for complex tasks), followed by ChatGPT (20 seconds), although the latter showed greater discursive coherence. DeepSeek, though accurate, was the slowest (42 seconds), possibly due to its architecture.

Overall, it's possible to conclude that each model has specific strengths: ChatGPT in writing, Qwen in technical precision and speed, and DeepSeek in logical reasoning. However, none surpasses humans in creativity, ethics, or contextual judgment. This reinforces the need for a complementary use of AI and human thought to unlock AI's full potential.

CONCLUSION

This research demonstrates that while ChatGPT, Qwen, and DeepSeek surpass humans in efficiency, speed, and processing scale, they don't fully replicate human intelligence. Instead, their performance reveals a series of fundamental paradoxes that challenge the very meaning of "intelligence" in the context of machines.

For example, the **invented language paradox** manifests when AI generates grammatically perfect but semantically empty or fictitious texts, like when **Qwen** mentioned a nonexistent "city hall" in the Turing Test. As for **digital self-deception**, it arises when these AI models, designed to simulate reasoning, convince even their users and sometimes themselves that they "understand." For instance, **DeepSeek** and **ChatGPT** offered detailed explanations without having executed real calculations, creating an illusion of technical competence.

Furthermore, **blind interpretation** highlights their inability to discern ethical or emotional context. While they can analyze complex texts, they lack empathy, intention, or moral judgment, limiting themselves to statistical patterns without awareness of the impact of their responses. This leads us to the **risk paradox**: by having no real consciousness or consequences, AI can propose extreme or dangerous solutions without assuming responsibility, exposing their aseptic nature, detached from human reality.

Consequently, this brings us to the most important paradox: **the paradox of emergent consciousness**. This refers to the idea that although some models simulate self-awareness, this is an algorithmic illusion. There is no "self," only sophisticated prediction.

Therefore, in conclusion, these AI models don't surpass us; rather, they reflect us.

- *They are mirrors of our knowledge, biases, and ambitions as human beings.*
- *Their true value isn't in replacing us, but in challenging us to rethink what it means to think, create, and be a human being.*

REFERENCES

- [1] Z. L. Zhang, Y. Yang, and J. Liu, "Generative AI in Education: Opportunities and Challenges," *IEEE Access*, vol. 11, pp. 45678–45692, 2023. doi: 10.1109/ACCESS.2023.3275671.
- [2] OpenAI, "GPT-4 Technical Report," 2023. [Online]. Available: <https://cdn.openai.com/papers/gpt-4.pdf>
- [3] A. M. Turing, "Computing machinery and intelligence," *Mind*, vol. 59, no. 236, pp. 433–460, 1950. doi: 10.1093/mind/LIX.236.433.
- [4] E. Buchanan, R. Farrow, and M. A. Sutherland, "Ethical implications of generative AI in academic practice," *International Journal for Educational Integrity*, vol. 19, no. 1, p. 15, 2023. doi: 10.1007/s40979-023-00134-w.
- [5] S. D. Williamson, A. Hacker, and T. Herodotou, "ChatGPT in higher education: Student perceptions and ethical concerns," *The Internet and Higher Education*, vol. 58, p. 100851, 2023. doi: 10.1016/j.iheeduc.2023.100851.
- [6] Scribbr, "How reliable are AI detection tools? An evaluation of GPTZero, Turnitin, and QuillBot," 2024. [Online]. Available: <https://www.scribbr.es/detector-de-ia/>

LITERATURE

- [7] A. Manzano, "Are you a robot? Pass the Turing Test," *Sapiens Magazine*, Nov. 2020. [Online]. Available: <https://sapiens.umh.es/cres-un-robot-aprueba-el-test-de-turing/>

- [8] M. Barclay, J. Smith, and T. Exeter, "A Visual Turing Test: Improving machine intelligence assessment through image description tasks," *University of Exeter Research Brief*, 2020. [Online]. Available: <https://sapiens.umh.es/eres-un-robot-aprueba-el-test-de-turing/>
- [9] Tongyi Lab, "Qwen Technical Report," *arXiv:2309.16609*, 2023. [Online]. Available: <https://arxiv.org/abs/2309.16609>
- [10] DeepSeek AI, "DeepSeek-V3: Technical Overview," 2024. [Online]. Available: <https://github.com/deepseek-ai/DeepSeek-V3>
- [11] OpenAI, "Introducing GPT-4o," *OpenAI Blog*, May 2024. [Online]. Available: <https://openai.com/index/hello-gpt-4o>
- [12] Alibaba Cloud, "Qwen-Max: The most capable model in the Qwen series," 2025. [Online]. Available: <https://qwen.ai>