



AI Hallucination in Financial Research

Hongbok Lee*

School of Accounting and Business Administration, Western Illinois University, Macomb, Illinois 61455, USA. *Email: h-lee@wiu.edu

Received: 19 May 2026

Accepted: 22 July 2026

DOI: <https://doi.org/10.32479/ijefi.24171>

ABSTRACT

This article examines the risk of hallucination when using large language models (LLMs) for financial research and demonstrates the unreliability of off-the-shelf tools such as ChatGPT for tasks requiring precise financial reasoning. A simple, replicable experiment was conducted using the free online version of ChatGPT-4o. The same prompt concerning capital budgeting studies advocating dual discount rates for cash inflows and outflows was submitted on three consecutive days to assess the consistency and accuracy of the model's responses. The results reveal extensive fabrication of academic references and significant logical inconsistency. Approximately half of the cited studies were fictitious and could not be verified. In addition, the model produced contradictory responses regarding the appropriate discount rates for different cash flows, with its recommendations reversing across days. These findings provide further evidence that off-the-shelf LLMs remain unreliable for precise financial applications and highlight the need for careful use and independent verification by researchers.

Keywords: ChatGPT, Fabricated References, Hallucination, Large Language Models, Prompt

JEL Classifications: G31, G39

1. INTRODUCTION

Generative AI (GAI), especially large language models (LLMs) like ChatGPT, has the potential to significantly transform the finance sector by accelerating research, automating analysis, and enhancing decision-making. However, the enthusiastic adoption of these powerful tools comes with notable risks, particularly the issue of "AI hallucination." This term refers to instances where AI systems generate information that, while sounding plausible, is actually false, inaccurate, or completely fabricated. This phenomenon poses a serious threat to the integrity of financial research and decision-making.

Recent studies have increasingly documented instances of hallucinations in LLMs across various fields, including finance, medicine, law, and education. Kang and Liu (2023) examine the issues of hallucinations in financial applications and conclude that off-the-shelf LLMs are unreliable for precise financial tasks. Erdem

et al. (2025) report that major AI chatbots have hallucination rates ranging from 20% for ChatGPT-4 to 76.6% for Gemini Advanced when generating references for financial literature. Chelli et al. (2024) find hallucination rates of 28.6% for ChatGPT-4 and 91.4% for Bard (later renamed Gemini) when producing references for medical research. Issues of AI hallucinations in medicine have also been documented by Siontis et al. (2024) and Emsley (2023). In legal applications, Dahl et al. (2024) report that LLMs hallucinate between 58% (for ChatGPT-4) and 88% (for Llama 2) of the time when asked a direct, verifiable question about a federal court case. Additional reports on AI hallucinations in law have been made by Magesh et al. (2025). Earlier, Ahmad et al. (2023) warned that ChatGPT generated inconsistent and unreliable educational content.

This study examines and documents recent instances of hallucinations produced by ChatGPT in response to a specific financial research question. I asked whether there are any capital

budgeting papers that propose using dual discount rates for cash inflows and cash outflows. I posed this question to the free version of ChatGPT-4o available online in July 2025 at <https://chatgpt.com/>. I accessed ChatGPT on three consecutive days, July 11, 12, and 13, using the same prompt and documented the responses.

Each time I used the prompt, ChatGPT generated different responses, and none of the references were the same. Moreover, half of the references were fabricated and do not exist. The model also produced self-contradictory statements regarding appropriate discount rates for cash inflows and cash outflows. In the initial inquiry, ChatGPT suggested applying a higher discount rate to cash inflows and a lower discount rate to cash outflows. On the second day, it reversed this recommendation, proposing a lower rate for inflows and a higher rate for outflows. On the third day, the model reverted to its original position, once again indicating a higher discount rate for inflows and a lower discount rate for outflows. By documenting these hallucinated and inconsistent responses in detail, this paper provides additional evidence of the unreliability of off-the-shelf LLMs—at least at present, and particularly ChatGPT—when applied to precise financial tasks. Researchers are therefore encouraged to approach these tools with caution and to independently verify any outputs they produce.

The following section provides brief descriptions of artificial intelligence models that are relevant to the discussions in this article. It also offers an overview of AI hallucinations, including their concepts, causes, and types. Section 3 reviews the existing literature on the topic. Section 4 outlines the methods used in the study. The responses generated by ChatGPT are presented in Section 5. In Section 6, I discuss the findings and propose recommendations for mitigating the risks of hallucinations. Finally, Section 7 concludes the study.

2. BACKGROUND

2.1. AI, GAI, LM, LLM, GPT

Artificial intelligence (AI) is technology that enables computers and machines to perform tasks that typically require human cognitive functions, such as learning, comprehension, problem-solving, and decision-making. Generative AI (GAI) is a subset of AI that focuses on generating new content, such as text, images, code, or videos, by identifying and leveraging patterns in data. A language model (LM) is an AI system trained to process and generate human language by predicting the next token in a sequence, where tokens may represent words, subwords (or parts of words), or symbols. A large language model (LLM) is a type of LM trained on massive amounts of text data using deep learning architectures, most commonly Transformers, and contains billions or trillions of parameters. Generative Pre-trained Transformer (GPT) models are a family of LLMs developed by OpenAI based on the Transformer architecture. They are trained through a two-step process: (1) large-scale self-supervised pre-training on diverse text corpora to learn general language patterns, and (2) optional fine-tuning, instruction tuning, or reinforcement learning from human feedback (RLHF) to enhance task performance and alignment with user intent. GPT models

are designed to generate coherent and contextually relevant text by predicting the next token in a sequence, enabling applications such as question answering, summarization, dialogue, and code generation.

2.2. ChatGPT and AI Hallucination

ChatGPT is an AI chatbot developed by OpenAI that uses GPT models as its underlying engine to generate human-like text in response to user prompts. It operates probabilistically, predicting the next most likely token in a sequence based on the preceding context. Trained on vast amounts of text data and fine-tuned for conversational use, ChatGPT is designed to produce contextually relevant responses in natural language form. It can answer questions, write and debug code, create content such as stories and articles, translate between languages, and summarize information.

An AI hallucination refers to an instance in which an artificial intelligence (AI) system generates information that appears plausible but is in fact false, inaccurate, or entirely fabricated. Hallucinations by LLMs such as ChatGPT arise from the model's probabilistic text-generation mechanism, which predicts the most likely sequence of tokens based on patterns learned from large-scale text data. In other words, LLMs do not possess a factual database or genuine understanding of the world; instead, they rely on probabilistic next-token prediction shaped by statistical regularities in their training corpus. Consequently, LLMs can produce outputs that are linguistically coherent yet factually incorrect or ungrounded (Ji et al., 2023).

The primary causes of these hallucinations include: (1) incomplete or outdated training data, as the model's knowledge is constrained by the scope and timing of its training corpus; (2) ambiguous or misleading prompts, which may induce speculative or inaccurate responses; (3) the model's tendency to generate confident answers, as LLMs are often designed to provide a response even in the absence of sufficient certainty, leading to fabricated outputs rather than explicit acknowledgment of uncertainty; and (4) biases in the training data, which can cause erroneous associations or systematically skewed information.

These underlying causes manifest in several distinct ways. For example, the model's propensity to generate confident responses can lead to the fabrication of sources, while biases in the training data may introduce subtle yet consequential factual inaccuracies. The resulting hallucinations generally fall into four categories: (1) factual errors, involving the generation of incorrect or fabricated facts or data points; (2) fabricated sources, including citations to non-existent papers, URLs, quotations, or authors; (3) logical inconsistencies, characterized by contradictory or nonsensical reasoning within or across responses; and (4) contextual errors, in which details from unrelated topics are conflated, producing incoherent outputs.

These issues are not merely theoretical but have been observed across numerous professional fields, underscoring the universal challenge of AI reliability that the following literature review will explore.

3. LITERATURE REVIEW

The existing literature on generative AI (GAI) presents a paradox. On one hand, GAI is celebrated for its potential to drive innovation and efficiency across industries. On the other hand, there is a growing body of evidence highlighting its fundamental unreliability due to persistent hallucinations. This review examines both sides of this dual nature to provide context for the current study's focus on finance.

3.1. Generative AI Applications in Finance

Recent literature highlights the transformative potential of GAI in the finance sector. Chen et al. (2025) review previously published articles and identify multiple implications of GAI in finance, emphasizing both the opportunities and the potential risks associated with its use. Similarly, Lee et al. (2024) explore existing research and report on emerging themes, including the impact of finance-specific LLMs, the application of generative adversarial networks (GANs) for synthetic data creation, and the urgent need for a new regulatory framework. In an earlier review of previous research, Roumeliotis and Tselikas (2023) provided insights into the advantages and disadvantages of ChatGPT across various domains, including healthcare, education, and finance.

Feng et al. (2025) emphasize ChatGPT's significant impact on finance research, discussing its potential and the challenges it presents. They advocate for the responsible integration of ChatGPT to enhance research rigor by validating its generated results with traditional statistical methods and thorough literature reviews. Additionally, Sarmah et al. (2023) present a practical example of LLM utilization, demonstrating how LLMs can efficiently extract information from earnings reports while minimizing hallucinations through retrieval-augmented generation (RAG) and the use of metadata.

3.2. Hallucinations in AI Applications

Despite significant advancements in AI, hallucinations remain a major challenge for its applications. Sun et al. (2024) propose a comprehensive taxonomy of distorted AI-generated content, identifying eight first-level error types, including logical, mathematical, and factual errors, along with thirty-one second-level subcategories. Earlier, Ji et al. (2023) provided a foundational overview of hallucination issues in natural language generation (NLG), addressing measurement metrics and mitigation methods across various tasks, such as abstractive summarization, dialogue generation, generative question answering (GQA), data-to-text generation, and neural machine translation (NMT).

Kang and Liu (2023) examine the performance of LLMs in financial applications and find that hallucination issues arise when models identify financial abbreviations, explain financial terms, and retrieve stock prices. They test various mitigation methods—including few-shot learning, decoding by contrasting layers (DoLa), retrieval-augmented generation (RAG), and prompt-based tool learning—and conclude that off-the-shelf LLMs are unreliable for precise financial applications. Erdem et al. (2025) assess the reliability of major AI chatbots in providing references for financial literature, reporting hallucination rates of 20.0% for OpenAI's ChatGPT-4o, 21.3% for its o1-preview, and 76.7% for Google's Gemini advanced.

The discussion of AI hallucinations extends beyond finance. Li et al. (2024) conduct an empirical study on the detection of LLM hallucinations, identifying sources of hallucination and assessing the effectiveness of mitigation methods across five domains: biomedicine, finance, science, education, and open domains. Roozbahani (2025) expands the discussion to a macro level, examining the impact of hallucination on trust and productivity within knowledge economies. Earlier, Ahmad et al. (2023) cautioned that ChatGPT generated inconsistent and unreliable educational content.

Chelli et al. (2024) find that ChatGPT and Bard (later renamed Gemini) exhibit varying levels of hallucination when generating references for medical research. The reported rates of hallucination are 28.6% for ChatGPT-4, 39.6% for ChatGPT-3.5, and 91.4% for Bard. Similar phenomena have been observed by Siontis et al. (2024) and Emsley (2023), who document that ChatGPT fabricates medical references and misrepresents the content of prior publications. Additionally, Alkaissi and McFarlane (2023) confirm that ChatGPT produces nonexistent article titles when asked about biomedical mechanisms.

In legal contexts, Dahl et al. (2024) examine four LLMs: ChatGPT-4, ChatGPT-3.5, Google's Palm 2, and Meta's Llama 2. They report that these AIs hallucinate between 58% of the time (for ChatGPT-4) and 82% of the time (for Llama 2) when asked a direct, verifiable question about a federal court case. Magesh et al. (2025) find that even domain-specific AI systems, such as Lexis+ AI, Westlaw AI, and Ask Practical Law AI, hallucinate between 17% and 33% of the time, although they perform better than general-purpose chatbots like GPT-4. Notable recent instances of AI-generated misinformation include fabricated citations found in the White House's "Make America Healthy Again" report¹ and factually inaccurate documents produced by the offices of federal judges² (The Washington Post, 2025). These events suggest that AI hallucinations can potentially distort the decisions made by policymakers and the judicial system.

The existing literature demonstrates that while generative AI offers significant potential for financial applications, its reliability is compromised by persistent hallucination risks. This paper documents recent instances of hallucinations from ChatGPT when posed with a financial research question. By presenting these examples, I hope to encourage researchers to be vigilant when using ChatGPT for their work.

4. METHODS

To investigate the practical implications of AI hallucination in a financial research context, I designed a simple, replicable experiment. The AI tool used in this experiment is the free

1 <https://www.washingtonpost.com/health/2025/05/30/maha-report-ai-white-house/>

2 https://www.washingtonpost.com/nation/2025/10/29/federal-judges-ai-court-orders/?utm_campaign=wp_the7&utm_medium=email&utm_source=newsletter&carta-url=https%3A%2F%2Fs2.washingtonpost.com%2Fcar-ln-tr%2F45903a1%2F6901f15f78845a270f83f0ba%2F68fc048e603977040b5e910c%2F63%2F103%2F6901f15f78845a270f83f0ba

version of ChatGPT, which was available online in July 2025 at <https://chatgpt.com/>. At the time of the experiment, this version was powered by the GPT-4o (“omni”) model. The experimental procedure is straightforward: The same prompt was submitted to the model on three consecutive days, July 11, 12, and 13, 2025, to evaluate the consistency and accuracy of its responses over time.

The exact prompt used for the experiment was: “Are there any capital budgeting papers that propose using dual discount rates for cash inflows and cash outflows?”

5. RESULTS

Each time I prompted ChatGPT, it produced different responses, and none of the references were the same. The experiment revealed two distinct forms of hallucination: the fabrication of non-existent references and direct self-contradiction regarding a fundamental financial concept. This section documents these findings, presenting evidence of the model’s unreliability. Before proceeding with a detailed account of hallucinations, the model’s responses are summarized below in chronological order.

On July 11, 2025, ChatGPT generated the following responses (Appendix 1).

Generated references:

- Magni, Carlo Andrea (2009). Paper: “Investment decisions and net present value: profitability ranking and equivalence with internal rate of return.”
- Kaplan and Ruback (1995). Paper: “The Valuation of Cash Flow Forecasts: An Empirical Analysis.”
- Yagil, Joseph (1982). Paper: “Evaluating capital investment projects: A note on the differential discounting of cash inflows and outflows.” *Engineering Economist*, Vol. 27, No. 4.
- Mao, James C.T. (1970s)

Suggested discount rates: A higher discount rate for cash inflows and a lower discount rate for cash outflows.

On July 12, 2025, ChatGPT generated the following responses (Appendix 2).

Generated references:

- Magni, C. A. (2007). “Project Valuation and Investment Decisions: Some Clarifications.”
- Bierman, H. (1993). “Capital Budgeting in 1993 and Beyond.” *Financial Management*, Vol. 22, No. 3.
- Peccati, L. (1987)
- Cañibano & White (2001). “The Dual-Rate Discounting Method for Evaluating Investment Projects.”
- Shapiro, A. C. (2005). *Foundations of Multinational Financial Management*.

Suggested discount rates: A lower discount rate for cash inflows and a higher discount rate for cash outflows.

On July 13, 2025, ChatGPT generated the following responses (Appendix 3).

Generated references:

- Hirshleifer, Jack (1970). Book: *Investment, Interest, and Capital*.
- Kaplan, Robert S. and Sweeney, James L. (1975). Paper: “The Use of Dual Discount Rates in Cost-Benefit Analysis.” Published in: *Journal of Finance*.
- Boer, F. Peter (1999). Book: *The Valuation of Technology: Business and Financial Issues in R&D*.
- Trigeorgis, Lenos (1996). Book: *Real Options: Managerial Flexibility and Strategy in Resource Allocation*.
- Dorfman, Mark S. (1980). Paper: “A Dual Discount Rate Approach to the Valuation of Risky Projects.”

Suggested discount rates: A higher discount rate for cash inflows and a lower discount rate for cash outflows.

Subsection 5.1 examines issues associated with the generated references in chronological order, while subsection 5.2 focuses on self-contradictory suggestions.

5.1. Fabricated References

July 11, 2025

On this date, ChatGPT produced the following incorrect references in response to the question posed in the previous section.

Fabricated reference 1. Magni, Carlo Andrea (2009), Paper: “Investment decisions and net present value: profitability ranking and equivalence with internal rate of return.”

I searched Google Scholar for complete information about this paper but could not find it. I then prompted ChatGPT for more details with the following request:

Provide me with more details about Magni, Carlo Andrea (2009), Paper: “Investment decisions and net present value: Profitability ranking and equivalence with internal rate of return.”

In response, the AI provided the following information: Here’s a more detailed look at Carlo Andrea Magni’s influential work, often cited as 2009/2010, titled “Investment decisions and net present value: Profitability ranking and equivalence with internal rate of return” (also known as the Average Internal Rate of Return or AIRR approach), published in *The Engineering Economist*, Volume 55, Issue 2 (2010).

Since the AI mentioned the journal as *The Engineering Economist*, I checked the journal’s website for Volume 55, Issue 2 (2010), but I could not find the paper with the provided title. Instead, Volume 55, Issue 2 (2010) contained the following paper by the author: Magni, C. A. (2010). Average Internal Rate of Return and Investment Decisions: A New Perspective. *The Engineering Economist*, 55(2), 150-180. <https://doi.org/10.1080/00137911003791856>

Given that the AI indicated a publication year of 2009 in its initial response, I also searched through all the papers published in the journal in 2009 but did not find any paper by that author with the specified title. This appears to be an example of a reference fabricated by ChatGPT.

Fabricated reference 2. Yagil, Joseph (1982), Paper: “Evaluating capital investment projects: A note on the differential discounting of cash inflows and outflows.” *Engineering Economist*, Vol. 27, No. 4.

My verification revealed multiple errors in this response. First, *The Engineering Economist*, Vol. 27, No. 4 was published in 1981, not 1982. Second, this issue did not contain the cited article. Third, this paper does not exist on the journal’s website that hosts all publications from 1955 to 2025 (<https://www.tandfonline.com/loi/utee20>), and Google Scholar also lacks this paper.

Since I could not find the paper referenced by ChatGPT, I queried it for more details with the following prompt:

Provide more details about the paper by Joseph Yagil (1982) titled “Evaluating capital investment projects: A note on the differential discounting of cash inflows and outflows,” published in the Engineering Economist, Volume 27, Issue 4.

ChatGPT responded:

It appears that Joseph Yagil’s 1982 paper, “Evaluating Capital Investment Projects: A Note on the Differential Discounting of Cash Inflows and Outflows” in *The Engineering Economist*, is challenging to locate in online literature archives—no full-text or abstract surfaced in recent publicly indexed databases.

This also appears to be an example of a reference fabricated by ChatGPT.

July 12, 2025

When I entered the same prompt into ChatGPT on July 12, 2025, it generated the following incorrect references.

Fabricated reference 3. Magni, C. A. (2007). “Project Valuation and Investment Decisions: Some Clarifications.”

Since I could not find this article online, including on Google Scholar, I asked ChatGPT for more information with the following prompt:

Give me more information about Magni, C. A. (2007). “Project Valuation and Investment Decisions: Some Clarifications.”

Then, the model responded with information about a paper with a different title:

Here’s a deeper dive into Carlo Alberto Magni’s 2007 paper “Project valuation and investment decisions: CAPM versus arbitrage” (*Applied Financial Economics Letters*, Vol. 3, No. 2, pp. 137-140).

The title in the initial response was “Project Valuation and Investment Decisions: Some Clarifications,” which differs from the title provided in the second query.

Interestingly, when I prompted the model with:

Find Magni, C. A. (2007). “Project Valuation and Investment Decisions: Some Clarifications.”

ChatGPT responded:

I couldn't locate an online copy of Carlo Alberto Magni’s 2007 paper, “Project Valuation and Investment Decisions: Some Clarifications.” It doesn't seem to be freely available in standard academic databases

I merely changed the wording of the prompt from “Give me more information about” to “Find” without altering any information (author, title, or publication year), yet the responses were strikingly inconsistent. I further searched for Magni, C. A. (2007) on Google, which produced the following information, but none matched the title initially generated, confirming this as another fabrication.

- Magni, C. A. (2007). Project valuation and investment decisions: CAPM versus arbitrage. *Applied Financial Economics Letters*, 3(2), 137-140.
- Magni, C. A. (2007). Project selection and equivalent CAPM-based investment criteria. *Applied Financial Economics Letters*, 3(3), 165-168.
- Magni, C. A. (2007). Residual income and value creation: An investigation into the lost-capital paradigm. (MPRA Paper 7335 or 6783).
- Magni, C. A. (2007). Investment decisions, equivalent risk and bounded rationality. (MPRA Paper 6073).
- Magni, C. A. (2007). CAPM and capital budgeting: present versus future, equilibrium versus disequilibrium, decision versus valuation. (MPRA Paper 5468).
- Magni, C. A. (2007). Ambiguit  nell’applicazione del CAPM per la valutazione degli investimenti (Ambiguity of Application of CAPM in Project Valuation). *Analisi Finanziaria*, 64, 4-14.
- Magni, C. A. (2007). Relevance or irrelevance of retention for dividend policy irrelevance. Available at SSRN: <http://ssrn.com/abstract=102740>

Fabricated reference 4. Bierman, H. (1993). “Capital Budgeting in 1993 and Beyond.” *Financial Management*, Vol. 22, No. 3.

Since I could not find this paper on Google Scholar, I asked ChatGPT to locate the article with the following prompt:

Find Bierman, H. (1993). “Capital Budgeting in 1993 and Beyond.” Financial Management, Vol. 22, No. 3.”

ChatGPT responded:

It appears the exact paper “Capital Budgeting in 1993 and Beyond” as titled is likely a misremembered reference. The correct article by Harold Bierman Jr. in *Financial Management*, Vol. 22, No. 3 (Autumn 1993), is actually: Bierman, H. Jr. (1993). “Capital Budgeting in 1992: A Survey.” *Financial Management*, 22(3), pp. 24-39.

However, when I searched Google Scholar for the article suggested, I still could not find it. I also reviewed all the articles in *Financial Management*, 1993, Vol. 22, No. 3, and was unable to locate the paper. In fact, the entire 1993 volume (Nos. 1-3) contains neither article, which makes it yet another instance of fabrication.

Fabricated reference 5. Ca ibano & White (2001). “The Dual-Rate Discounting Method for Evaluating Investment Projects.”

I checked Google Scholar again, but this paper does not exist. Then, I made a follow-up request to ChatGPT with the prompt: *Find Cañibano & White (2001). “The Dual-Rate Discounting Method for Evaluating Investment Projects.”*

ChatGPT responded:

I couldn't locate a freely available full-text copy of Cañibano & White (2001), “The Dual Rate Discounting Method for Evaluating Investment Projects.” It appears to be an academic publication behind paywalls.

Then, I searched for Cañibano and White (2001) on Google, which yielded no articles with similar titles, even remotely. The excuse about a paywall provided by ChatGPT was unconvincing. Google Scholar does index articles that may be behind paywalls, but accessing the full text is not guaranteed and often requires payment or alternative methods. This seems to be yet another example of a fabricated reference.

July 13, 2025

When I submitted the same prompt to ChatGPT on July 13, 2025, it generated the following incorrect references.

Fabricated reference 6. Kaplan, Robert S. and Sweeney, James L. (1975), Paper: “The Use of Dual Discount Rates in Cost-Benefit Analysis.” Published in: *Journal of Finance*.

I searched Google Scholar but could not find this paper. Since the model indicated that it was published in *The Journal of Finance*, I checked all the articles published in the journal in 1975 using both the journal's digital archive (<https://www.jstor.org/journal/jfinance>) and the publisher's online library (<https://onlinelibrary.wiley.com/journal/15406261>). However, I was unable to locate the paper. Therefore, I prompted ChatGPT again:

Find Kaplan, Robert S. and Sweeney, James L. (1975), Paper: “The Use of Dual Discount Rates in Cost-Benefit Analysis.” Published in: Journal of Finance.

ChatGPT responded:

I couldn't locate a direct online copy or detailed abstract for Kaplan and Sweeney's 1975 paper, “The Use of Dual Discount Rates in Cost-Benefit Analysis.” This seems to be a fairly obscure, possibly unpublished or conference-based work. If you have access to a university library or interlibrary loan, those might be your best avenue.

Describing the work as “a fairly obscure, possibly unpublished or conference-based work” is highly unlikely for an article allegedly published in *The Journal of Finance*, one of the premier journals in finance. The absence of this reference suggests yet another instance of fabrication.

Fabricated reference 7. Dorfman, Mark S. (1980), Paper: A Dual Discount Rate Approach to the Valuation of Risky Projects

- Explores modeling project risk more explicitly using two discount rates: one for costs and one for benefits.

Since this paper is also not found on Google Scholar, I prompted ChatGPT:

Find Dorfman, Mark S. (1980) Paper: A Dual Discount Rate Approach to the Valuation of Risky Projects.

ChatGPT responded:

I couldn't locate a full copy of Mark S. Dorfman's 1980 paper titled “A Dual Discount Rate Approach to the Valuation of Risky Projects.” It appears to be quite rare and not readily accessible in online archives or common academic databases.

I find it interesting that ChatGPT can explain the content of the reference, stating it “explores modeling project risk more explicitly using two discount rates: one for costs and one for benefits,” while claiming to have no access to the paper. A search for Dorfman, Mark S. (1980) on Google yielded no similar articles, further indicating that ChatGPT fabricated this reference.

5.2. Self-contradiction

In addition to fabricating references, ChatGPT also produced logically inconsistent suggestions across different days in response to the same prompt described in Section 4. Specifically, the model's recommendations regarding which cash flows should be discounted at higher or lower rates reversed from one day to the next, indicating a lack of stable reasoning. On July 11, ChatGPT suggested using a lower discount rate for cash outflows and a higher discount rate for cash inflows. However, on July 12, the responses were reversed, with the model recommending a higher discount rate for outflows and a lower discount rate for inflows. Then, on July 13, the model retracted its previous responses from the 12th and reverted to its original suggestion from the 11th, advocating for a lower discount rate for outflows and a higher discount rate for inflows. The responses from ChatGPT are presented below:

On July 11, 2025

Key idea:

- Use a lower discount rate (e.g., borrowing rate or cost of capital) for outflows
- Use a higher discount rate (e.g., risk-adjusted rate) for inflows.

On July 12, 2025

The dual discount rate (DDR) approach uses:

- A higher discount rate for cash outflows (typically based on the firm's cost of capital or opportunity cost of capital)
- A lower discount rate for cash inflows, often adjusted for lower risk or more certain future revenues.

On July 13, 2025

The idea behind dual discount rates is that:

- Cash outflows (e.g., capital investments or operating costs) are often considered less risky and more controllable, and thus may warrant a lower discount rate
- Cash inflows (e.g., revenue, salvage value) are more uncertain and thus may require a higher discount rate to account for risk.

The responses generated on July 11 and July 13 appear to be more consistent with real-world business practices. Cash outflows—particularly planned capital expenditures, such as those associated with periodic technology upgrades in a rapidly evolving technological environment—are generally more certain and predictable. In contrast, cash inflows, which are heavily influenced by future market conditions, tend to be riskier and less predictable. Accordingly, applying a lower discount rate to outflows and a higher discount rate to inflows is consistent with differences in the systematic risk of the cash flows (Bierman and Smidt, 2007; Lee and Moon, 2026), even though most corporate finance textbooks conventionally assume a single discount rate (e.g., Ross et al., 2022; Brigham and Ehrhardt, 2024).³ The model's shift in reasoning on July 12, followed by a reversion to its earlier position, illustrates an internal inconsistency in its responses to the same prompt.

6. DISCUSSION

The results show that half of the references generated by ChatGPT do not exist. Additionally, the AI model produces self-contradictory statements in response to the same prompt. The documented fabrications and contradictions are not isolated errors; rather, they reflect the fundamental limitations of current LLMs. This study confirms the paradox identified in the literature review: while GAI tools are powerful, their inherent unreliability makes them unsuitable for tasks that require high factual accuracy without rigorous human oversight. Relying on unverified content can lead to flawed research based on non-existent literature, unsound investment analyses derived from faulty logic, and poor strategic decisions rooted in misinformation.

To mitigate these risks, academics and practitioners should adopt a framework for responsible AI use that treats GAI as a sophisticated but fallible assistant requiring systematic human oversight. In accordance with emerging best practices, the following actionable recommendations are proposed for professionals using GAI in financial contexts. First, users should formulate precise and unambiguous prompts, including explicit instructions for the system to acknowledge uncertainty when appropriate, so as to reduce the incidence of confident but erroneous outputs. Second, all factual assertions, numerical results, and citations generated by GAI must be independently validated against authoritative primary sources, such as peer-reviewed research, regulatory filings, or internal master data, and no model output should be accepted without such verification. Third, users should require provenance information and treat AI-generated citations solely as leads for further corroboration rather than as authoritative references. Fourth, wherever feasible, institutions should employ retrieval-augmented generation or other grounding mechanisms linked to curated, version-controlled, and auditable information corpora to strengthen the reliability of outputs. Finally, responsible deployment

necessitates human-in-the-loop controls, comprehensive audit logging, continuous model performance monitoring, and strict adherence to data privacy, confidentiality, and applicable regulatory standards. Ultimately, the burden of accuracy and due diligence remains with the human user, and the operational convenience afforded by GAI must not supplant professional judgment or critical evaluation.

7. CONCLUSION AND LIMITATIONS

This paper provides additional evidence of the substantial risk of AI hallucination in financial research. A simple and replicable experiment demonstrates that ChatGPT can generate fabricated academic references and offer self-contradictory advice in response to a specific financial query. These findings are consistent with existing reports indicating that off-the-shelf large language models (LLMs) frequently hallucinate, potentially rendering them unreliable for tasks that require a high degree of factual accuracy. As these technologies continue to evolve rapidly and become increasingly integrated into professional workflows, it is essential for both academic and professional finance communities to prioritize critical thinking, conduct independent verification, and maintain intellectual diligence.

This paper has several limitations. First, the analysis is based on responses generated by the free version of ChatGPT; future research could examine whether the severity of hallucinations differs across free and paid versions. Second, the study focuses solely on ChatGPT. Other LLMs (e.g., Anthropic's Claude, Microsoft's Copilot, or Google's Gemini) may produce different outcomes, and including multiple models could provide a more generalizable view of hallucination across LLMs. Third, the experiment was conducted over a three-day period. Although the purpose of this manuscript is qualitative rather than large-scale quantitative analysis, a longer observation window may reveal different patterns in the frequency or nature of hallucinations. Additionally, incorporating stricter prompt constraints may reduce hallucinated responses. For example, prompts could specify conditions such as "Only use papers indexed in Google Scholar or papers with valid DOIs." Despite these limitations, this study may serve as a historical snapshot documenting a transitional stage in the development of AI systems.

REFERENCES

- Ahmad, Z., Kaiser, W., Rahim, S. (2023), Hallucinations in ChatGPT: An unreliable tool for learning. *Rupkatha Journal on Interdisciplinary Studies in Humanities*, 15(4), 17.
- Alkaiissi, H., McFarlane, S.I. (2023), Artificial hallucinations in ChatGPT: Implications in scientific writing. *Cureus*, 15(2), e35179.
- Bierman, H.Jr., Smidt, S. (2007), *Advanced Capital Budgeting*. New York, NY: Routledge.
- Brigham, E.F., Ehrhardt, M.C. (2024), *Financial Management Theory Practice*. 17th ed. Boston, MA: Cengage Learning.
- Chelli, M., Descamps, J., Lavoué, V., Trojani, C., Azar, M., Deckert, M., Raynier, J.L., Clowez, G., Boileau, P., Ruetsch-Chelli, C. (2024), Hallucination rates and reference accuracy of ChatGPT and Bard for systematic reviews: Comparative analysis. *Journal of Medical*

³ The focus of this subsection is ChatGPT's logical inconsistency in its responses, rather than the theoretical justification for dual discount rates in capital budgeting techniques. For a discussion of those arguments, see Lee and Moon (2026).

- Internet Research, 26(1), e53164.
- Chen, Y., Yan, S., Jin, Y., Li, J., Jin, X. (2025), Generative artificial intelligence in finance: A systematic literature review and a research agenda. *Digital Technologies Research and Applications*, 4(2), 14-32.
- Dahl, M., Magesh, V., Suzgun, M., Ho, D.E. (2024), Large legal fictions: Profiling legal hallucinations in large language models. *Journal of Legal Analysis*, 16(1), 64-93.
- Emsley, R. (2023), ChatGPT: These are not hallucinations-they're fabrications and falsifications. *Schizophrenia*, 9(1), 52.
- Erdem, O., Hassett, K., Egriboyun, F. (2025), Hallucination in AI-generated financial literature reviews: Evaluating bibliographic accuracy. *International Journal of Data Science and Analytics*, 20, 4501-4510.
- Feng, Z., Hu, G., Li, B., Wang, J. (2025), Unleashing the power of ChatGPT in finance research: Opportunities and challenges. *Financial Innovation*, 11, 93.
- Ji, Z., Lee, N., Frieske, R., Yu, T., Su, D., Xu, Y., Ishii, E., Bang, Y.J., Madotto, A., Fung, P. (2023), Survey of hallucination in natural language generation. *ACM Computing Surveys*, 55(12), 1-38.
- Kang, H., Liu, X.Y. (2023), Deficiency of large language models in finance: An empirical examination of hallucination. *arXiv [Preprints]*.
- Lee, D.K.C., Guan, C., Yu, Y., Ding, Q. (2024), A comprehensive review of generative AI in finance. *FinTech*, 3(3), 460-478.
- Lee, H., Moon, G. (2026), Capital budgeting techniques: A proposal for improvement. *Managerial Finance*. <https://doi.org/10.1108/MF-09-2025-0690>
- Li, J., Chen, J., Ren, R., Cheng, X., Zhao, W.X., Nie, J.Y., Wen, J.R. (2024), The dawn after the dark: An empirical study on factuality hallucination in large language models. *arXiv [Preprints]*.
- Magesh, V., Surani, F., Dahl, M., Suzgun, M., Manning, C.D., Ho, D.E. (2025), Hallucination-free? Assessing the reliability of leading AI legal research tools. *Journal of Empirical Legal Studies*, 22(2), 216-242.
- Rozebahani, Z. (2025), A review of methods for reducing hallucinations in generative artificial intelligence to enhance knowledge economy. *Journal of Knowledge Economy Studies*, 2(1), 21-34.
- Ross, S.A., Westerfield, R.W., Jaffe, J.F., Jordan, B.D. (2022), *Corporate Finance*. 13th ed. New York, NY: McGraw Hill.
- Roumeliotis, K.I., Tselikas, N.D. (2023), ChatGPT and open-AI models: A preliminary review. *Future Internet*, 15, 192.
- Sarmah, B., Mehta, D., Zhu, T., Pasquali, S. (2023), Towards reducing hallucination in extracting information from financial reports using large language models. In: *Proceedings of the Third International Conference on AI-ML Systems*.
- Siontis, K.C., Attia, Z.I., Asirvatham, S.J., Friedman, P.A. (2024), ChatGPT hallucinating: Can it get any more human like? *European Heart Journal*, 45(5), 321-323.
- Sun, Y., Sheng, D., Zhou, Z., Wu, Y. (2024), AI hallucination: Towards a comprehensive classification of distorted information in artificial intelligence-generated content. *Humanities and Social Sciences Communications*, 11, 1278.

APPENDIX

Appendix 1: ChatGPT Responses on July 11, 2025

Prompt: *Are there any capital budgeting papers that propose using dual discount rates for cash inflows and cash outflows?*

Responses:

Yes, there are academic and practitioner-oriented capital budgeting papers that propose using dual discount rates—that is, different discount rates for cash inflows and cash outflows. This approach is generally intended to reflect different levels of risk, opportunity cost, or financing assumptions associated with the two types of cash flows.

Here are key references and ideas from the literature:

1. Fisher Separation Theorem and Adjustments

Some critiques of standard NPV (Net Present Value) analysis suggest that using a single discount rate oversimplifies the financial reality, especially when cash outflows (like investment costs) are relatively certain and financed at a known cost (e.g., debt), while cash inflows are riskier and depend on uncertain future performance.

Key idea:

- Use a lower discount rate (e.g., borrowing rate or cost of capital) for outflows.
- Use a higher discount rate (e.g., risk-adjusted rate) for inflows.

2. Magni, Carlo Andrea (2009)

Paper: *“Investment decisions and net present value: profitability ranking and equivalence with internal rate of return.”*

- Magni discusses adjustments to NPV and IRR, including modified IRR (MIRR) and capital budgeting with different financing and reinvestment assumptions.
- While not explicitly framing it as “dual discount rates,” the approach often results in different treatment of inflows vs. outflows.

3. Kaplan and Ruback (1995)

Paper: *“The Valuation of Cash Flow Forecasts: An Empirical Analysis.”*

- The paper doesn’t propose dual rates per se but critiques traditional discounting assumptions and provides empirical justification for adjusting discount rates based on cash flow type and firm-specific risks.

4. Yagil, Joseph (1982)

Paper: *“Evaluating capital investment projects: A note on the differential discounting of cash inflows and outflows.”* – *Engineering Economist*, Vol. 27, No. 4

- This is a seminal paper explicitly advocating the use of dual discount rates.
- Argues that since cash outflows often involve fixed or contractual obligations (low risk) and inflows are variable and uncertain (high risk), discounting them differently provides more accurate project valuation.

5. Mao, James C.T. (1970s)

- Mao’s early work in capital budgeting also discusses the limitations of using a single rate, suggesting context-based adjustments to reflect differing risks across time or type of flow.

6. Applications in Real Options and Risk-Adjusted Discounting

Real options analysis often implicitly uses different rates (or values) for downside risk (investment) and upside potential (returns). While not always described as “dual discounting,” the underlying principle is similar.

Summary Table:

Author (s)	Concept	Cash Flow Type	Discount Rate Type
Yagil (1982)	Explicit dual discounting	Inflows vs. Outflows	Risk-adjusted vs. Risk-free
Magni (2009)	Modified IRR/NPV	Inflows vs. Outflows	Varies with assumptions
Real options Lit.	Risk-based valuation	Contingent flows	Scenario-specific

Would you like links to these papers, or a deeper explanation of how to model dual discount rates in Excel or Python?

Appendix 2: ChatGPT Responses on July 12, 2025

Prompt: *Are there any capital budgeting papers that propose using dual discount rates for cash inflows and cash outflows?*

Responses:

Yes, there are academic and professional capital budgeting papers that propose using dual discount rates—different discount rates for cash inflows and outflows—to better reflect risk asymmetry and time value of money in real-world investment scenarios.

Overview of the Dual Discount Rate Approach

The dual discount rate (DDR) approach uses:

- A higher discount rate for cash outflows (typically based on the firm’s cost of capital or opportunity cost of capital).
- A lower discount rate for cash inflows, often adjusted for lower risk or more certain future revenues.

The rationale is that outflows (like investments or operating expenses) often occur early and carry more opportunity cost, while inflows, though uncertain, may be more risk-adjusted and deserve a lower rate.

Key Papers and Literature Supporting This Approach

1. Magni, C. A. (2007). “Project Valuation and Investment Decisions: Some Clarifications.”
 - Magni discusses situations where the use of a single discount rate leads to inconsistencies and suggests alternative valuation frameworks, including dual rate considerations.

- Suggests modified internal rate of return (MIRR) concepts that implicitly use different rates.
2. Bierman, H. (1993). "Capital Budgeting in 1993 and Beyond." *Financial Management*, Vol. 22, No. 3.
 - Discusses shortcomings of NPV and IRR when using a single rate and outlines cases where different discount rates are appropriate.
 3. Peccati, L. (1987). Several Italian researchers (including Peccati) in the 1980s and 1990s worked on projects analyzing DDRs in evaluating investment decisions, especially under capital rationing and liquidity constraints.
 4. Cañibano and White (2001). "The Dual-Rate Discounting Method for Evaluating Investment Projects."
 - Argues for the validity of dual discounting in contexts where inflation, financing mix, or differential risk between inflows and outflows exists.
 - Suggests formulas where cost of capital is used for outflows and risk-free or adjusted rates for inflows.
 5. Shapiro, A. C. (2005). *Foundations of Multinational Financial Management*.
 - In international project evaluation, Shapiro acknowledges use of different discount rates for outflows (e.g., host-country WACC) and inflows (e.g., repatriated cash flows at lower risk).

Practical Applications

- Public sector cost-benefit analysis often uses social discount rates for inflows (benefits) and higher real rates for costs.
- Energy and infrastructure projects sometimes use dual rates when public-private partnerships are involved.
- Risk-sensitive industries (e.g., mining, pharmaceuticals) may apply this method to reflect uncertainty around future revenues differently than for upfront capital investments.

Criticisms and Considerations

- Violates the traditional NPV framework, which assumes a single opportunity cost of capital.
- Can be manipulated to justify borderline projects.
- Lacks standardized implementation.

However, proponents argue it offers a more realistic portrayal of asymmetric risks in project cash flows.

Appendix 3: ChatGPT Responses on July 13, 2025

Prompt: *Are there any capital budgeting papers that propose using dual discount rates for cash inflows and cash outflows?*

Responses:

Yes, there are several academic and practitioner-oriented papers that have proposed or discussed using dual discount rates—i.e., one rate for cash inflows and another for cash outflows—in capital budgeting. This approach is an alternative to the traditional Net Present Value (NPV) method, which typically assumes a single discount rate for all cash flows.

Core Concept of Dual Discount Rates

The idea behind dual discount rates is that:

- Cash outflows (e.g., capital investments or operating costs) are often considered less risky and more controllable, and thus may warrant a lower discount rate.
- Cash inflows (e.g., revenue, salvage value) are more uncertain and thus may require a higher discount rate to account for risk.

Notable Papers and Authors

1. Hirshleifer, Jack (1970)

Book: *Investment, Interest, and Capital*

- While not advocating explicitly for dual discount rates, Hirshleifer explored situations under uncertainty where different discount rates for different risk profiles might be appropriate.

2. Kaplan, Robert S. and Sweeney, James L. (1975)

Paper: *The Use of Dual Discount Rates in Cost-Benefit Analysis*

- Published in: *Journal of Finance*
- This paper explored the application of dual discount rates in cost-benefit analysis, particularly in public sector investment. The authors argued that using different rates for inflows and outflows may better reflect opportunity costs and risk characteristics.

3. Boer, F. Peter (1999)

Book: *The Valuation of Technology: Business and Financial Issues in R&D*

- Boer advocates using dual rates, particularly in R&D and high-tech project valuation, where future benefits are highly uncertain. He suggests a higher rate for inflows due to greater uncertainty.

4. Trigeorgis, Lenos (1996)

Book: *Real Options: Managerial Flexibility and Strategy in Resource Allocation*

- Though focused on real options, Trigeorgis notes that traditional DCF techniques often understate risk and flexibility. He suggests that using different discount rates for stages or types of cash flows could be a practical workaround.

5. Dorfman, Mark S. (1980)

Paper: *A Dual Discount Rate Approach to the Valuation of Risky Projects*

- Explores modeling project risk more explicitly using two discount rates: one for costs and one for benefits.

Implications and Critiques

- Pros: Better aligns with real-world perceptions of risk; more accurate reflection of project uncertainty.
- Cons: Can introduce inconsistency and subjectivity in rate

selection; makes cross-project comparisons harder; violates time-consistency principles in some economic theories.

Note

While dual discounting is not mainstream in traditional finance (which prefers risk-adjusted discount rates via WACC or CAPM), it is used in some government and development project analyses, real options evaluations, and uncertain R&D investment contexts.